



Introduction to vLLM: High-Performance LLM Serving & Maximizing Performance on Limited GPUs

vLLM Thailand Meetup - Sept 2026

Tun Jian Tan ([@tjanaa](#))

vLLM Maintainer

Principal AI Engineer, Embedded LLM

Slides co-prepared by Kuan-fu Liu, Embedded LLM





VISION

vLLM everywhere:
the fastest engine
on any hardware,
powering an interoperable,
federated AI utility.

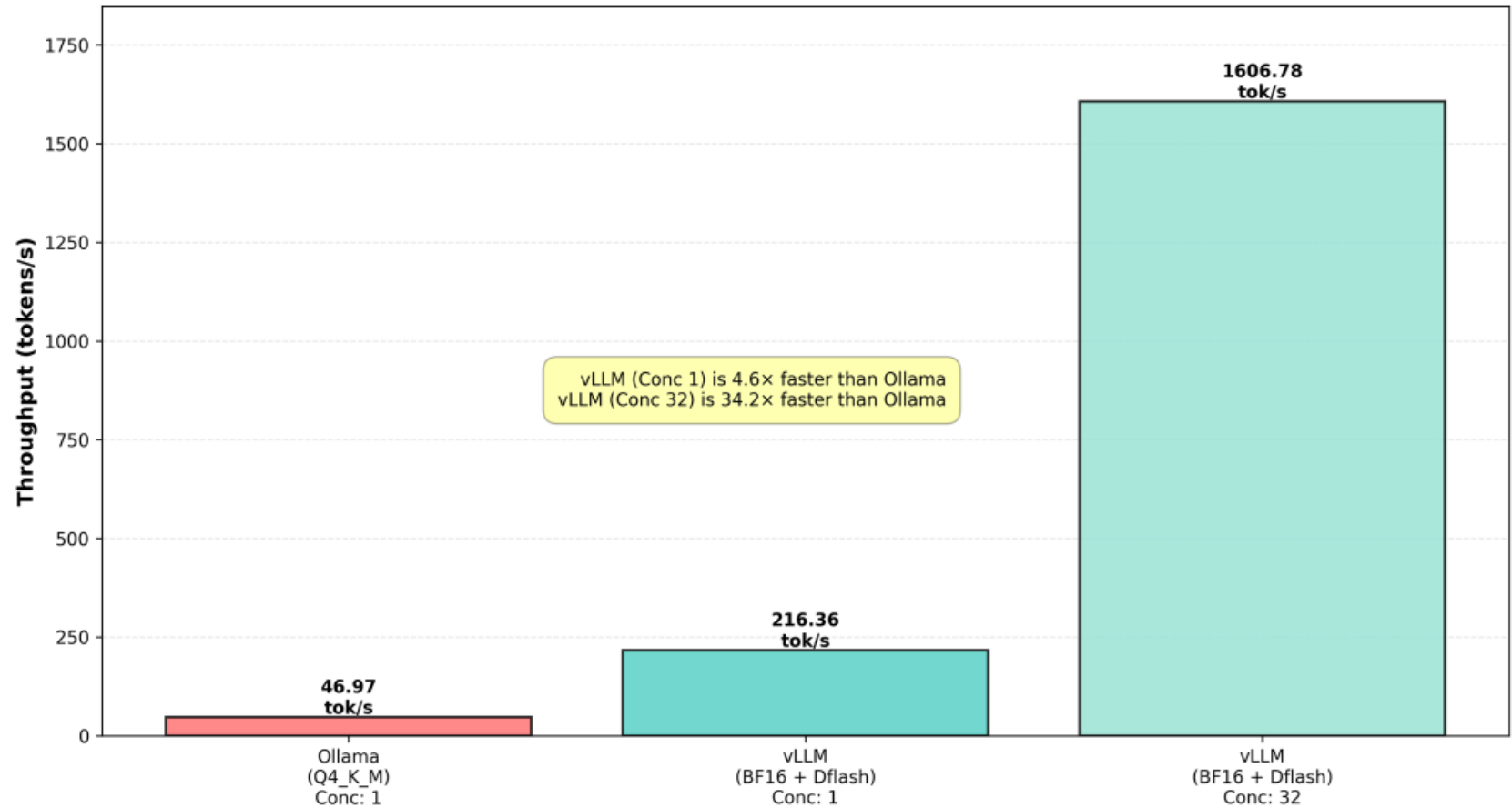
MISSION

A federated AI utility powered by
vLLM: fastest on any hardware,
interoperable by design, and
aligned with the ecosystem so
contributors and partners all
benefit.

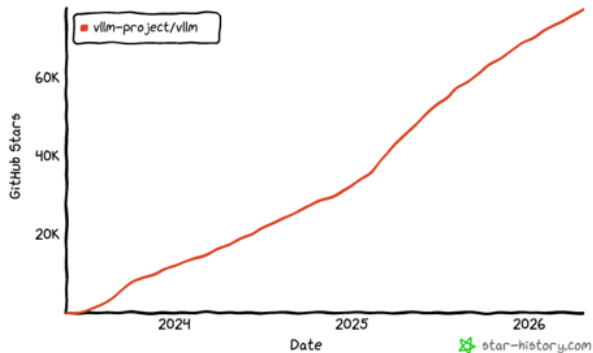
Different category of software

	Ollama	vLLM
Target	Single user, laptop	Concurrent, production
Batching	1 (batch size by default)	Continuous batching
Scale-out	TP (dependent on platform)	TP/PP/DP/EP, PD Disaggregated Inferencing
Day-0 support	Lagging	Same day
Quantization	GGUF	FP8, INT4, INT8, MXFP4, NVFP4

Throughput Comparison: Ollama vs vLLM (Different Concurrency Levels)



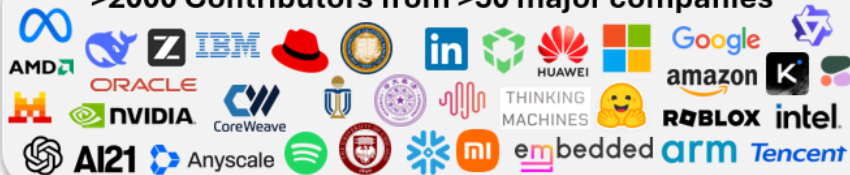
Star History



Most Popular LLM Serving Engine

- **91K+** GitHub stars, **800+** PRs/month
- **500K++** GPUs deployed 24/7
- **12.5K+** members in slack.vllm.ai

>2000 Contributors from >50 major companies



Broad Model Support (>100 arches)



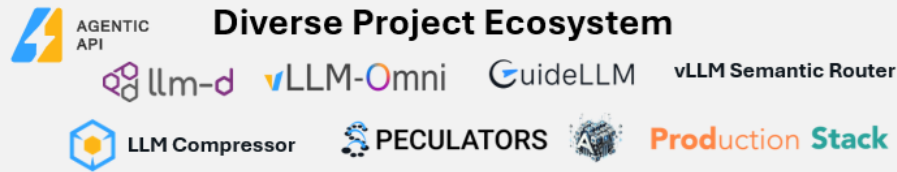
Flexible Device Parallelism

Tensor, Pipeline, Expert, Data, Context Parallel, Disagg Prefill/Decode, Disagg Encoder

Wide Hardware Support



Diverse Project Ecosystem





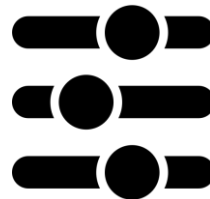
Security

Your data never leaves your VPC. Sensitive prompts stay inside your trust boundary



Privacy

No vendor telemetry on your prompts. No mystery retention policy. No black box.



Control

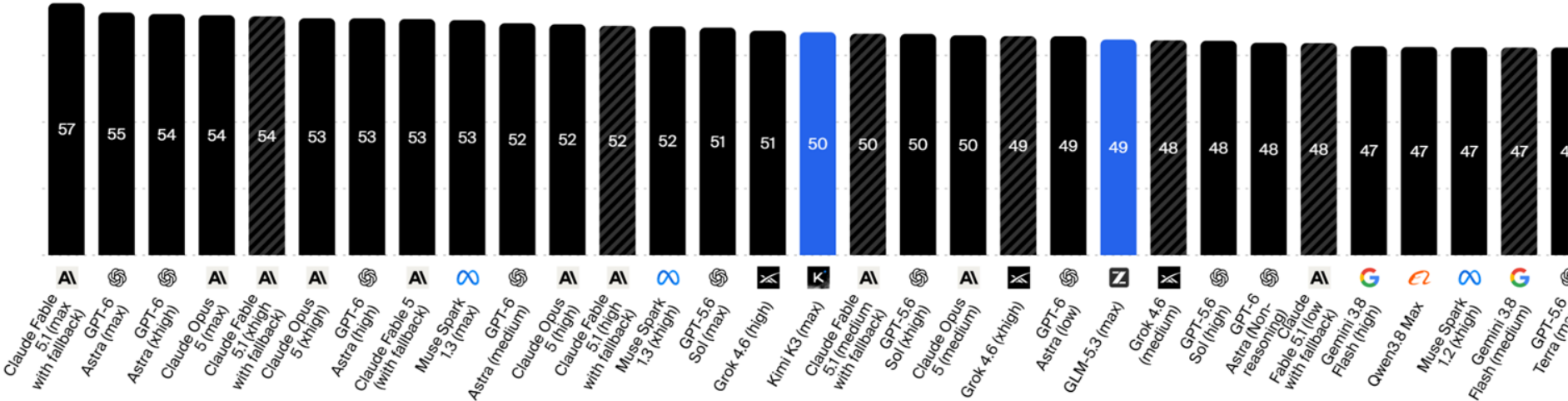
Your model. Your routing. Your SLAs. Your own pace on upgrades and policy.

Artificial Analysis Intelligence Index by Open Weights / Proprietary

Artificial Analysis Intelligence Index v4.2 incorporates 10 evaluations: AA-Briefcase, GDPval-AA v2, r³-Banking, Terminal-Bench v2.1, SciCode, Humanity's Last Exam, GDP.pdf, CritPt, AA-Omniscience, AA-LCR v1.1

- Estimate (independent evaluation forthcoming)
- Proprietary ● Open Weights (Commercial Use Restricted) ● Open Weights ● Open Weights (Non-commercial)

Artificial Analysis

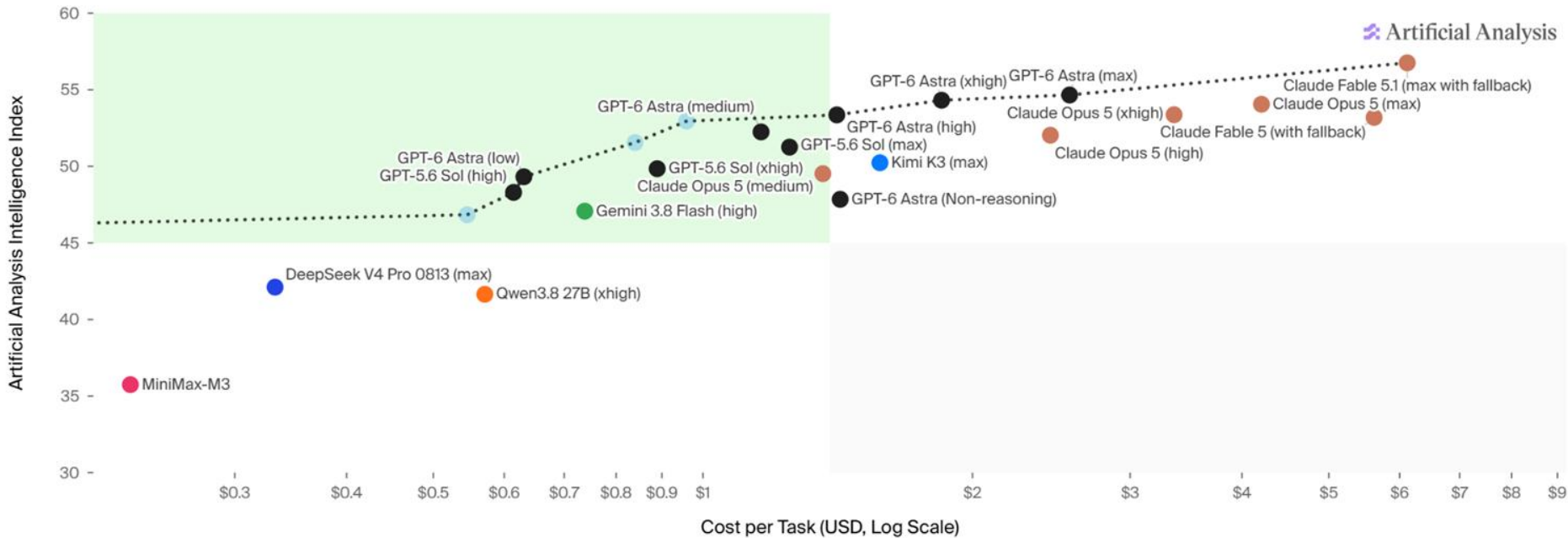


Intelligence Index vs. Cost per Intelligence Index Task

Artificial Analysis Intelligence Index · Weighted average cost (USD) per Artificial Analysis Intelligence Index task

Most attractive quadrant Pareto line

Anthropic OpenAI Alibaba Kimi MiniMax Google DeepSeek Meta Z AI Xiaomi InclusionAI IBM



Consumer GPUs vs Datacenter GPUs

CONSUMER GPUs



DATACENTER GPUs

H100 (80GB)

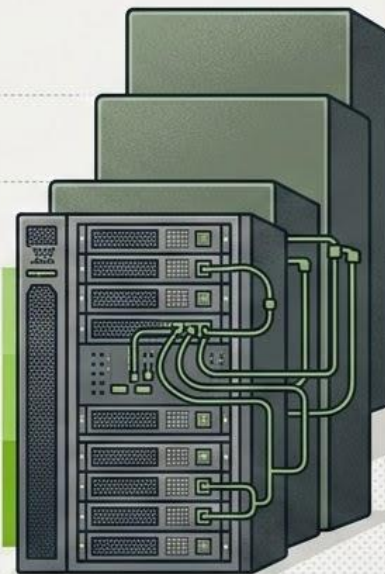
B300 (288GB)

B200 (192 GB)

H200 (141GB)

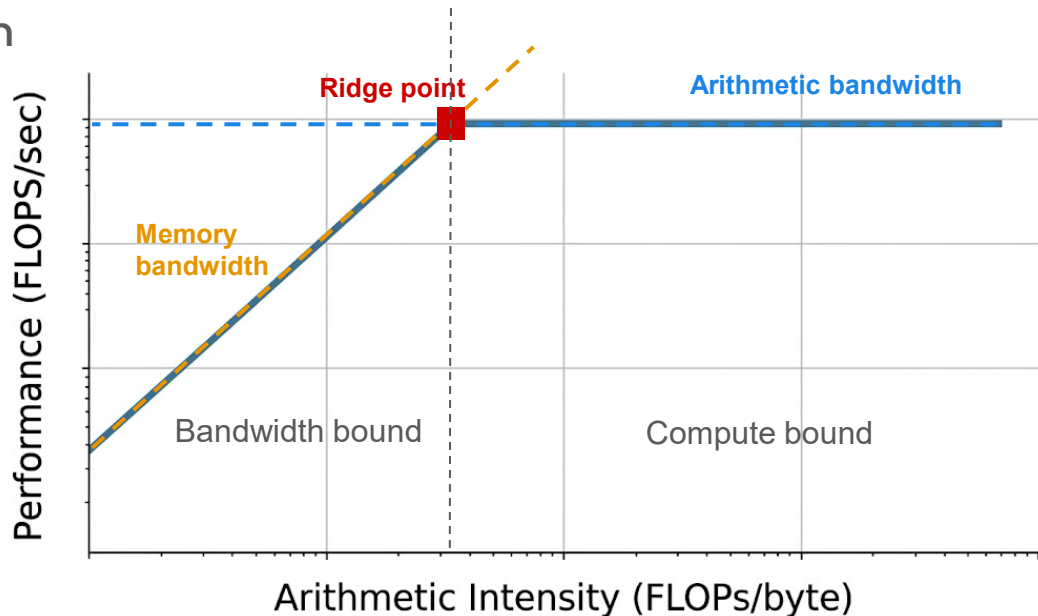
GB300 (288GB)

H100 (80GB)
VRAM



Hardware constraints in vLLM serving

- Compute power
- Memory bandwidth
- Memory space



What is consuming VRAM

- **Model weights**

$$Memory_{weight} = N_{parameters} * Bytes_{weight}$$

- e.g. Llama-3.1-70B with BF16 weight: 141GB

- **KV cache**

$$Bytes_{token} = 2 * L * H_{kv} * d_{head} * Bytes_{kv}$$

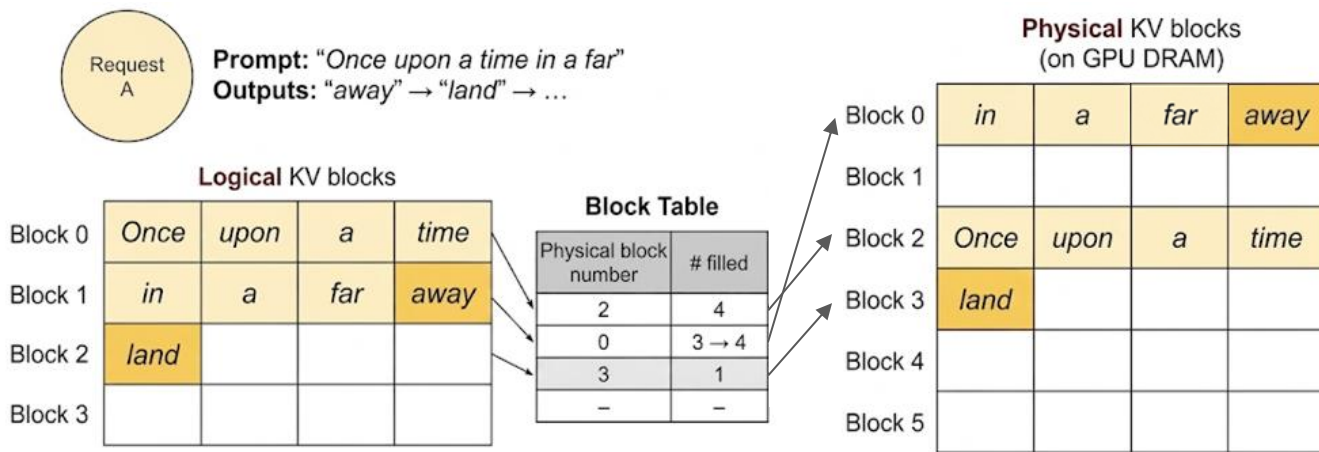
- e.g. Llama-3.1-70B with BF16 KV cache: 2.56GB with 8k sequence
- SSM or hybrid models

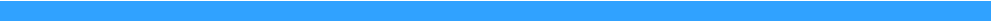
- **Static memory**

- CUDA graph
- Pytorch memory pools
- Kernel workspace buffers

KV Cache

- Skip costly KV recalculation
- Paged attention divides sequences into blocks
- LRU eviction policy





How to make the most of your GPUs

- **Quantization**
 - Model weights
 - Activations
 - KV cache
- **Improve cache hits**
 - Attend to prefix caches
- **Workload optimization**
 - Balance compute and memory bandwidth

Model weights/activations quantization

$$Memory_{weight} = N_{parameters} * Bytes_{weight}$$

- BF16 -> FP8: saves half the space
- vLLM supports quantized model weights
 - GGUF, GPTQ, AWQ, llm-compressor, bitsandbytes, etc
- Online quantization at model loading
 - `--quantization [fp8|...]`
- Faster hardware tensor/matrix arithmetics for quantized workloads

Radeon™ RX 9070 XT

Peak Half Precision (FP16 Vector) Performance 48.7 TFLOPs

Peak Half Precision (FP16 Matrix) Performance 195 TFLOPs

Peak Half Precision (FP16 Matrix) Performance with Structured Sparsity 389 TFLOPs

Peak 8-bit Precision (FP8 Matrix) Performance (E5M2, E4M3) 389 TFLOPs

Peak 8-bit Precision (FP8 Matrix) Performance with Structured Sparsity (E5M2, E4M3) 779 TFLOPs

Peak 8-bit Precision (INT8 Matrix) Performance 389 TOPs

Peak 8-bit Precision (INT8 Matrix) Performance with Structured Sparsity 779 TOPs

KV cache quantization

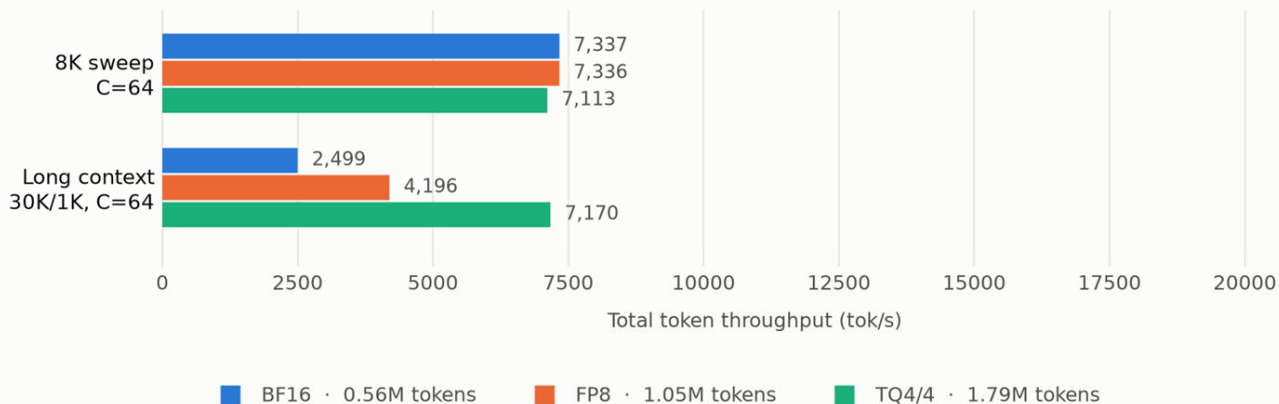
$$Bytes_{token} = 2 * L * H_{kv} * d_{head} * Bytes_{kv}$$

- Store more tokens in memory
 - `--kv-cache-dtype fp8`
 - `--calculate-kv-scales` # Calibration
 - Also supports calibration with external datasets
- Less memory traffic for bandwidth bound workloads

KV cache quantization

KV cache dtype vs throughput — Qwen3.5-27B

2x MI355X, TP=2, KV budget pinned at 19.7 GB/GPU. Capacity per dtype is in the legend.



Capacity shown is from the throughput servers; the agentic servers ran with prefix caching on, which costs 6-18% of the pool (0.52M, 0.93M, 1.47M respectively).

Automatic prefix caching

- Re-use KV cache across requests
 - `--enable-prefix-caching`

Request A

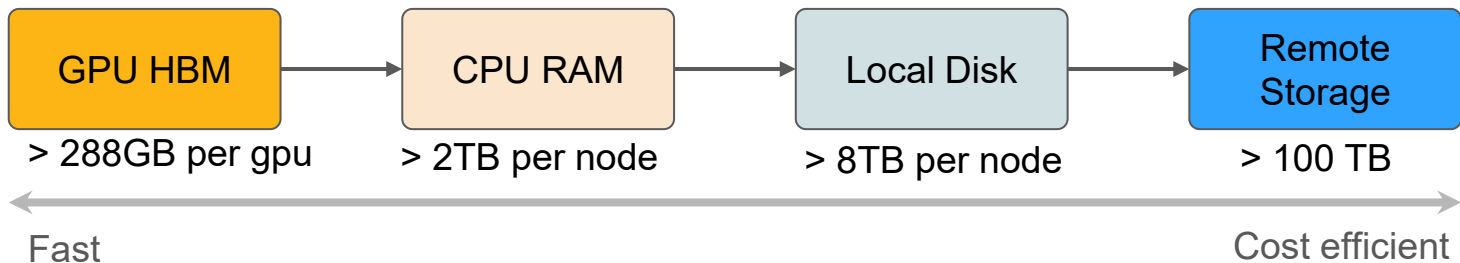
You are a helpful assistant gives detailed and polite answers to the user's questions. User. Hello!

Request B

You are a helpful assistant gives detailed and polite answers to the user's questions. User. How's the weather?

KV cache offloading

- In multi-turn agentic workloads, context can get very long and re-computing kv cache can be expensive
- In combination with **LMCache**, vLLM's KV cache can be offloaded to CPU DRAM, NVME, or remote storage

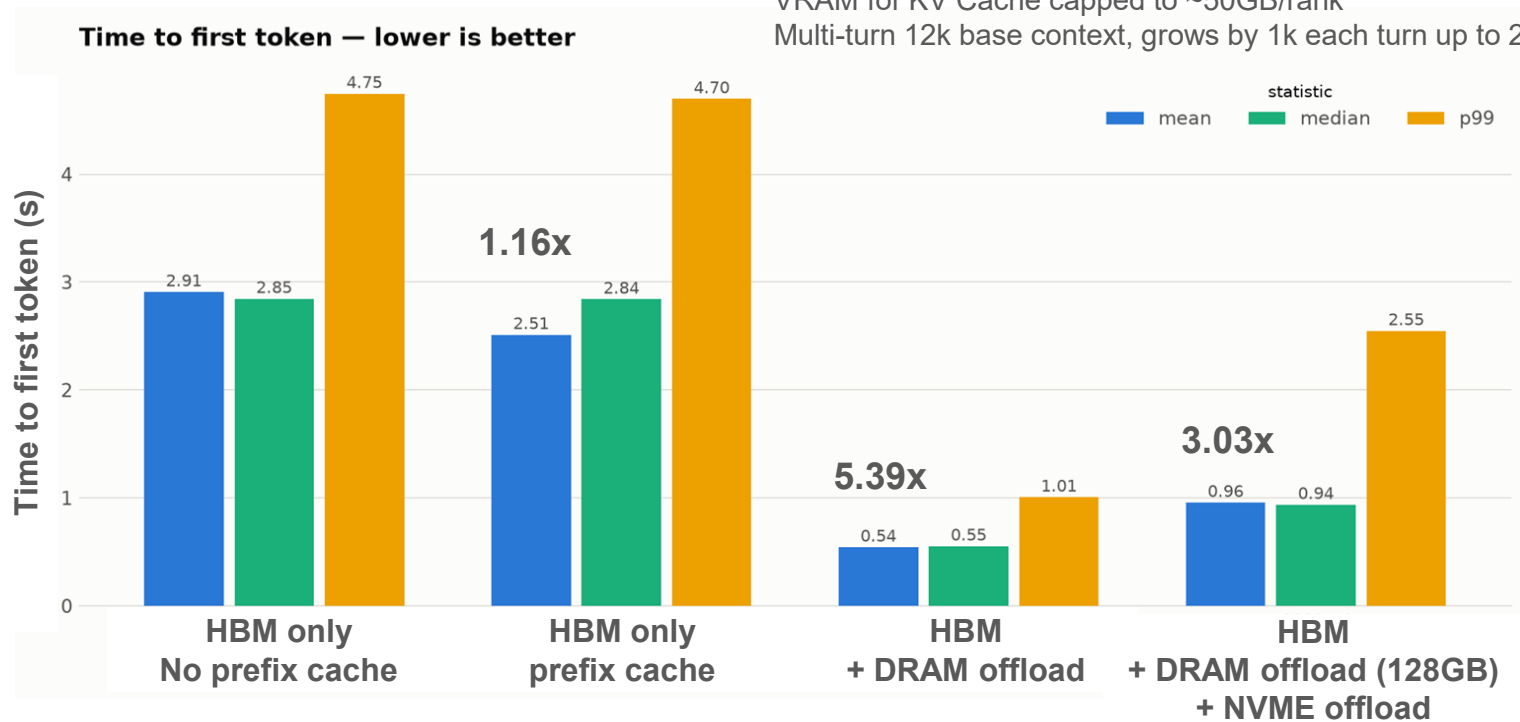


KV cache offloading

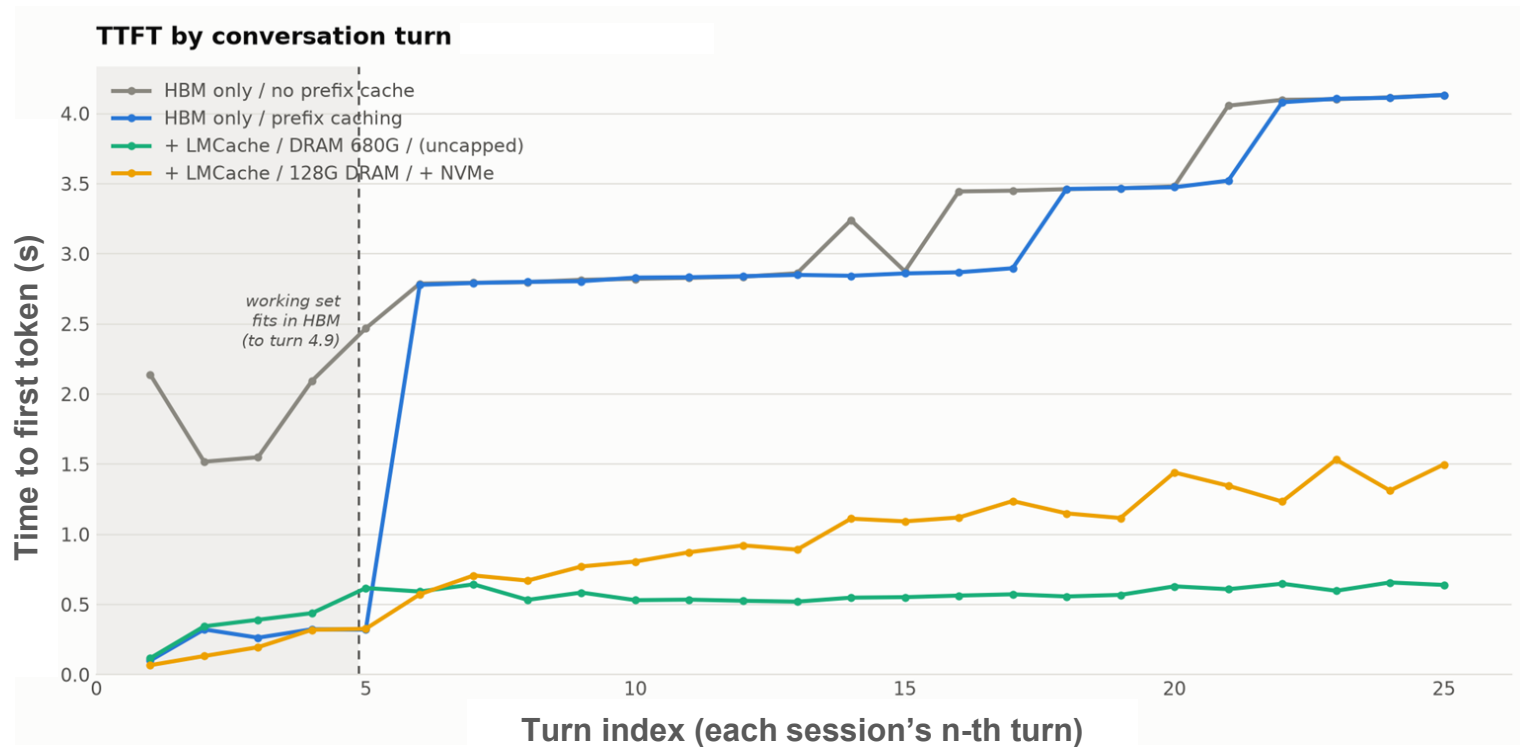
MiniMax-M3 TP4 on MI355X

VRAM for KV Cache capped to ~50GB/rank

Multi-turn 12k base context, grows by 1k each turn up to 25 turns, C=16



KV cache offloading



Cost Tier: 3 Year Rental

Updated: 2026-09-04

Source: SemiAnalysis InferenceX™

TCO \$/chip/hr:

MI355X: 2.1

Source: SemiAnalysis Market July 2026 Pricing Surveys & AI Cloud TCO Model [🔗](#)

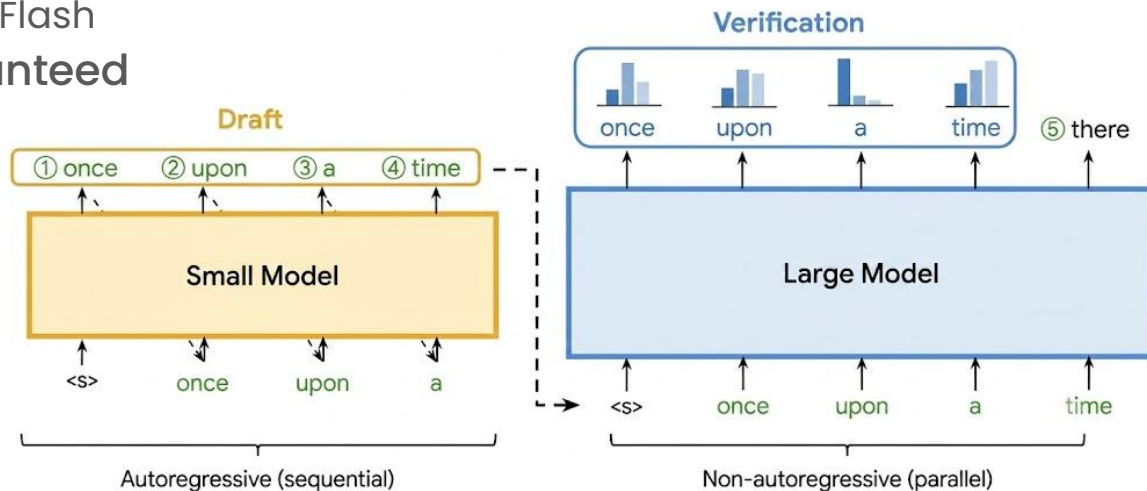


Quick Filters



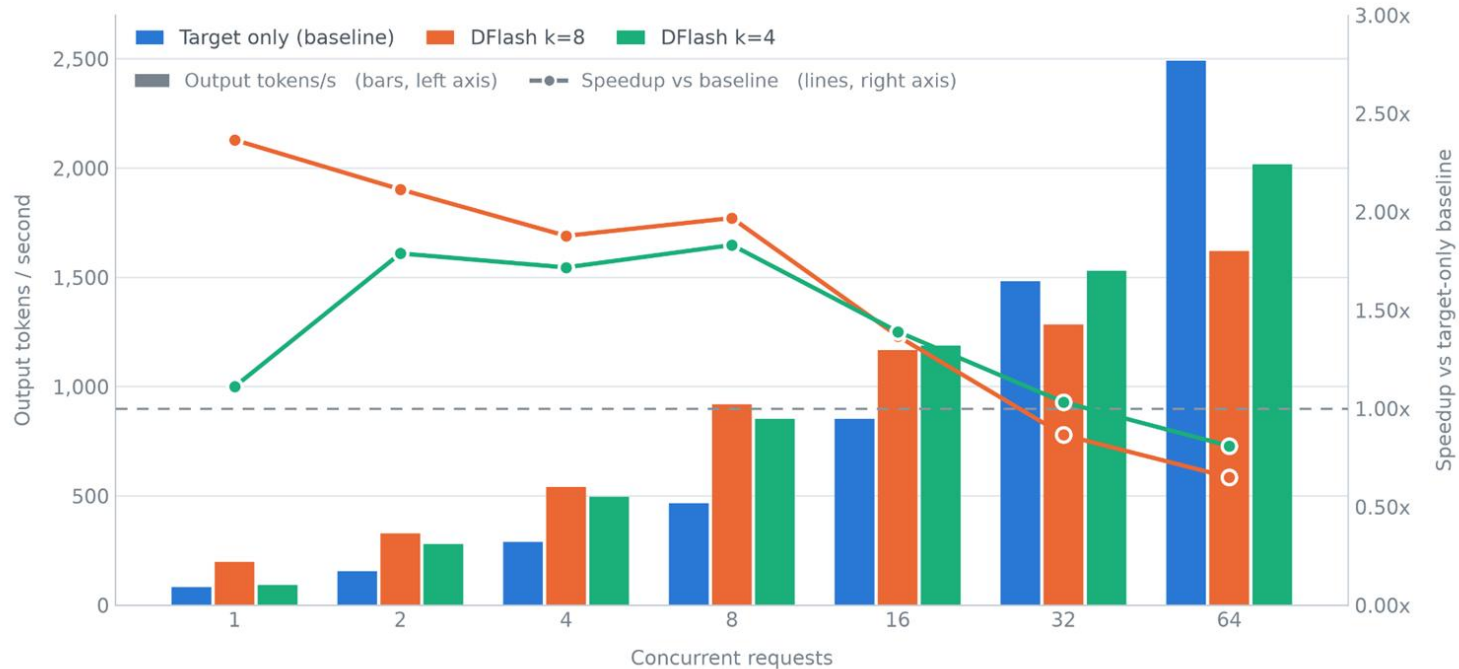
Speculative Decoding

- **Reduce ITL in low-medium concurrency workloads**
 - Pushing workload from memory bound regime more towards compute
- **Separate draft model or weights**
 - MTP, EAGLE, DFlash
- **Accuracy guaranteed**



Speculative Decoding

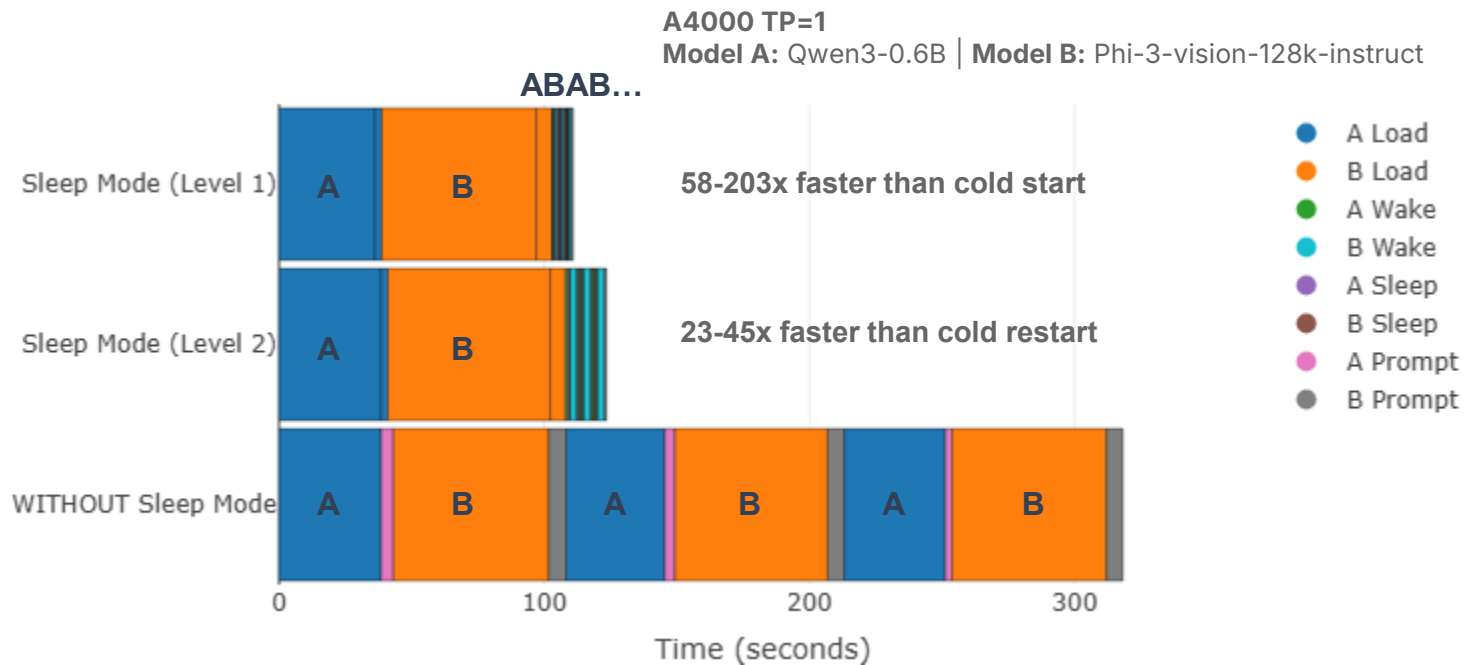
Qwen3.5-27B TP=2 on MI355
with z-lab/Qwen3.5-27B-DFlash

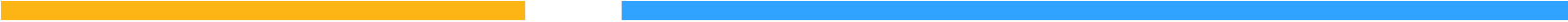


Sleep mode

- Reducing cold-start time for multi-model switching on the same hardware
- **Level 1: Weight Offloading** – Offloads weights to CPU RAM, discards KV cache
 - Memory allocator, process state, CUDA graphs, and compiled kernels retained
 - Provides the fastest wake times
 - Best for frequent model switching when CPU memory is sufficient
- **Level 2: Memory Optimization** – Discards both weights and KV cache from CPU RAM
 - Best for limited CPU RAM that can't hold all served models
 - Avoids reinitialization overhead, still much faster than a cold restart

Sleep mode





They are not free gains

- **Quantization – potential accuracy drop**
 - Be sure to verify with evaluation tools (lm-eval, lmms-eval, etc)
- **Memory hardware**
 - Additional memory spaces and devices
- **Draft models and MTP layers**
 - Require additional VRAM
 - Pushing batching and speculative decoding together can over-saturate compute

What to keep in mind in setting up vLLM

Know your workload

Before allocating system resources, analyze the demands of your specific deployment environment.

- **Concurrency Requirements**

Identify how many concurrent users or requests the system must serve.

- **Task Characteristics**

Determine the types of tasks (e.g., short chat inputs, massive document generation, structured extraction).

Choose strategies based on workloads

Tailor your model configurations and software optimizations to match your analyzed query patterns.

- **Optimization Techniques**

Select parameters such as chunked prefill, tensor parallelism, or speculative decoding based on traffic density.

- **Resource Allocation**

Balance compute budget against memory limits to maximize throughput without exceeding hardware bounds.

Get Involved with vLLM

Building the fastest and easiest-to-use open-source LLM inference & serving engine together!



<https://vllm.ai/>

Explore vLLM and join Weekly SIG meetings



<https://github.com/vllm-project/vllm>

Star, file issues, send PRs



<https://slack.vllm.ai>

12.5K+ members and SIG channels



<https://www.linkedin.com/company/vllm-project>

Follow and connect to get latest news



https://twitter.com/vllm_project

Follow and connect to get latest news



<https://roadmap.vllm.ai>

Follow the Q2 board

If you are serving LLMs in
2026,
you're either using 📄LLM or
explaining why you're not.

Thank you
-embedded