
LLMVisor: A Real-Time Latency Attribution Model for Multi-Tenant LLM Serving

Shuowei Jin^{1*} Xueshen Liu^{1*} Jiaxin Shan² Le Xu²
Tieying Zhang² Liguang Xie² Z. Morley Mao¹

¹University of Michigan ²ByteDance

Abstract

As LLM inference shifts to multi-tenant GPU clusters, co-batching improves throughput but obscures per-tenant usage and limits control. Enabling fractional sharing of the inference engine requires a real-time, per-request attribution primitive that is accurate and light enough to run inside the scheduling loop. We present *LLMVisor*, a roofline-guided *latency attribution* model that captures the memory-bound and compute-bound phases via a concise piecewise-linear form over features proportional to FLOPs and memory I/O traffic. *LLMVisor* decomposes batch latency into additive, per-request shares and runs efficiently at μ s-scale. We evaluate *LLMVisor* across Llama3.1-8B and Qwen2.5-14B/32B on A100/H100 under varying tensor parallelism and workload mixes. Compared to a token-count proxy baseline, *LLMVisor* attains near-perfect R^2 and reduces relative error by up to $2.5 \times / 3.3 \times$ (p90/p99) for prefill and $3.5 \times / 4.4 \times$ for decode, despite batching variability and sequence divergence.

1 Introduction

Large language models (LLMs) now power search, coding assistants, and conversational systems at scale [4, 7, 5]. To keep accelerators busy, providers co-batch heterogeneous requests via shared public endpoints [15, 2]. While batching boosts throughput, it leaves tenants with opaque, provider-tuned scheduling and little leverage to allocate or account for their own usage. The common alternative—renting dedicated GPU servers—squanders parallel hardware on interactive, small-batch workloads; for example, GPT-J 6B on A100 achieves only 0.4% SM utilization at batch size 1 and remains low even for 1,024-token sequences [9].

We argue for virtualizing the *inference engine*: tenants reserve fractional shares of GPU time and submit requests into shared batches, akin to CPU isolation via VMs, cgroups, and containers [3, 11, 12], and to credit/burst offerings (e.g., AWS $\tau 3$.*) [1]. Hardware partitioning (e.g., NVIDIA MIG) provides strong isolation but is inflexible for LLM memory footprints [14]. Realizing a software path to fractional sharing hinges on a missing primitive: a *real-time, per-request latency attribution model* accurate enough to guide scheduling and light enough to run in the critical path without harming batching efficiency. Such attribution directly enables SLO-aware admission (curbing tails by rejecting or deferring requests whose shares would exceed budgets), transparent post-hoc accounting, and fast “what-if” planning (e.g., alternative batch sizes or mixes) while preserving the benefits of co-batching.

Designing this primitive is difficult for several reasons. First, co-batching couples requests: their costs are not trivially separable. Second, LLM inference has phase-specific bottlenecks, compute-bound *prefill* versus memory-bound *decode*, that require different treatments. Third, costs scale nonlinearly with sequence length (e.g., quadratic self-attention) and depend on context length (KV-cache loads) and batch size. Fourth, throughput shifts with model architecture, tensor-parallel layout, and hardware. Finally, the scheduler operates at millisecond granularity, so attribution must run at μ s scale to be

*Equal contribution.

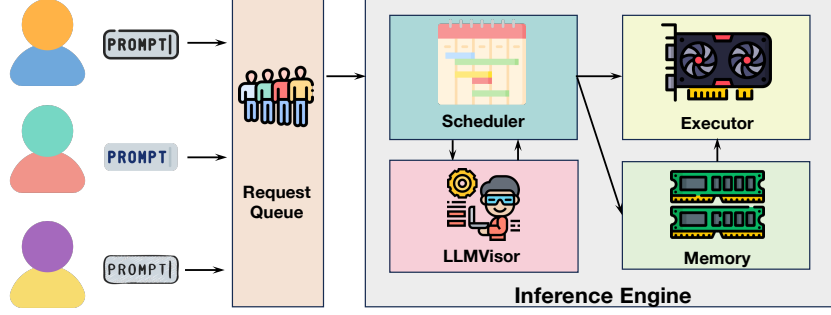


Figure 1: A typical workflow in modern large language model (LLM) inference engines can be summarized as follows. Upon arrival, each request is enqueued in a request buffer. At each inference step, the scheduler selects requests from the queue for execution. *LLMVisor* operates during the scheduling stage by providing GPU time attribution, enabling the scheduler to enforce multi-tenant fairness according to reserved GPU time quotas. Once scheduled, the inference engine dispatches the requests to the GPU, where the corresponding KV cache and model weights are loaded to perform the inference computation.

37 viable. Black-box ML predictors can fit latency trends but lack additivity and are too slow for
 38 the loop [18]; naive token-count proxies [16] ignore context and batching effects, yielding poor
 39 high-percentile accuracy.

40 We present *LLMVisor*, a roofline-guided *latency attribution* model for multi-tenant LLM serving.
 41 *LLMVisor* captures prefill and decode with a concise piecewise-linear form over features proportional
 42 to FLOPs and memory I/O traffic—explicitly encoding quadratic self-attention, context-length-driven
 43 KV-cache traffic, and batch-size effects—so end-to-end batch latency decomposes into additive, per-
 44 request shares. The resulting structure supports online μ s-scale “what-if” queries for admission
 45 control and budgeting while preserving batching efficiency. A lightweight prototype integrated into
 46 vLLM runs in the scheduling path with negligible overhead and achieves low error across diverse
 47 models and GPUs.

48 Our contributions are threefold:

- 49 • **Inference-aware modeling:** An interpretable, roofline-guided, piecewise-linear formulation that
 50 embeds LLM execution mechanics: compute-bound prefill vs. memory-bound decode, quadratic
 51 self-attention, KV-cache traffic via context length, and batch-size utilization effects. The model is
 52 additive by design, yielding closed-form per-request latency shares suitable for multi-tenant control.
- 53 • **Efficiency:** Microsecond-level runtime per scheduling step (100x faster than common ML predic-
 54 tors), enabling deployment *inside* the scheduler without disrupting batching.
- 55 • **Accuracy and generality:** Across Llama3.1-8B and Qwen2.5-14B/32B on A100/H100 with
 56 varying tensor parallelism and workload mixes, *LLMVisor* attains near-perfect R^2 and reduces
 57 relative error by up to $2.5 \times / 3.3 \times$ (p90/p99) for prefill and $3.5 \times / 4.4 \times$ for decode.

58 2 *LLMVisor* Design

59 **Design requirements.** The state-of-the-art inference engine workflow is illustrated in Fig. 1. *LLMVi-*
 60 *visor* assists the scheduler in ensuring that each tenant receives its reserved share of GPU time. To be
 61 practical in high-frequency scheduling loops (e.g., every decode step), *LLMVisor* must combine three
 62 properties: it must be *accurate*, providing reliable predictions of end-to-end batch latency; it must
 63 support *attribution fidelity*, decomposing that latency into per-request and per-tenant contributions;
 64 and it must be *efficient*, introducing negligible overhead so that scheduling can operate at each step².
 65 *LLMVisor* is designed to satisfy all three simultaneously.

66 **Problem formulation.** Formally, let $\mathcal{B} = \{1, \dots, |\mathcal{B}|\}$ denote a batch of requests and $u(i)$ the tenant
 67 of request i . We seek a model f that (i) predicts the end-to-end latency of the batch, $f(\mathcal{B}) \approx T_{\mathcal{B}}$,

²Step is the minimal execution unit in LLM inference. During decoding, one step corresponds to all active requests in the batch generating a single token. During prefill, one step corresponds to all requests in the batch completing their entire prompt processing.

and (ii) assigns each request a nonnegative share $f(i \mid \mathcal{B}) \approx t_i$ such that $\sum_{i \in \mathcal{B}} f(i \mid \mathcal{B}) = f(\mathcal{B})$. These per-request attributions capture each request’s latency contribution and aggregate naturally to per-tenant usage $L_u = \sum_{i: u(i)=u} f(i \mid \mathcal{B})$, supporting fairness, admission control, and billing.

Related work. Prior studies on LLM latency prediction follow two main approaches. The first uses ML predictors [6], which can be accurate but cannot attribute latency to individual requests, limiting their usefulness for multi-tenant scheduling. The second relies on analytical, roofline-style modeling [18], classifying inference as compute- or memory-bound and estimating latency from FLOPs and memory traffic relative to GPU peaks. While interpretable, these models are often inaccurate in practice due to modeling assumptions and time-varying GPU throughput, and they do not provide lightweight, per-request attribution in real time. VTC [16] attributes usage by token counts, which is imprecise because compute and I/O vary with token position [8].

2.1 Piecewise roofline-guided linear latency model

LLMVisor builds on the intuition of roofline analysis but adapts it for higher accuracy and real-time use. Rather than hand-computing FLOPs and memory bytes from model architecture and dividing by peak GPU rates, we construct a small set of request-level features proportional to FLOPs and I/O traffic and fit their coefficients using short warm-up profiling runs. This preserves interpretability while improving accuracy and keeping runtime cost negligible.

Let p_i denote the number of tokens *processed* for request i in the current step (all prompt tokens during prefill; $p_i = 1$ during decode), and let c_i denote the *context* tokens that must be attended and whose KV cache must be loaded (zero in prefill; prior prompt length in decode). With batch size $|\mathcal{B}|$, we model latency as a two-segment linear function corresponding to the two bottleneck phases:

$$T_{\mathcal{B}} = \begin{cases} \beta_{\text{pre}} + a_1 \sum_i p_i + a_2 \sum_i c_i + a_3 \sum_i p_i^2 + a_4 \sum_i |\mathcal{B}|, & \text{Compute-bound,} \\ \beta_{\text{dec}} + d_1 \sum_i p_i + d_2 \sum_i c_i + d_3 \sum_i p_i^2 + d_4 \sum_i |\mathcal{B}|, & \text{Memory-bound.} \end{cases} \quad (1)$$

Here, $\beta_{\star}$ captures fixed per-batch costs (e.g., model weight I/O); $\sum_i p_i$ relates to MLP and causal attention compute; $\sum_i c_i$ captures KV-cache memory traffic; $\sum_i p_i^2$ models quadratic self-attention costs; and $|\mathcal{B}|$ captures utilization gain. Coefficients (a, d, β) are fit via ordinary least squares on the profiling data.

Because $T_{\mathcal{B}}$ is linear in sums of per-request terms, *LLMVisor* naturally admits **closed-form**, additive per-request attributions via Eq. 1. yielding each request’s estimated latency contribution in the critical path.

Additionally, in terms of **efficiency**, Eq. 1 is over 100x faster than ML predictors such as Random Forest. In our measurements, *LLMVisor* operates at the microsecond level per request, whereas more complex ML models like Random Forest require milliseconds for inference. This efficiency makes *LLMVisor* well-suited for integration into LLM inference engine scheduling, which runs at millisecond granularity.

3 Evaluation

In this section, we evaluate *LLMVisor* across a diverse range of real-world LLM serving configurations and compare its predictions against ground truth measurements.

3.1 Experimental Setup

We evaluate *LLMVisor* inside vLLM (v0.7.3) [10] using its internal profiler for per-step ground truth, spanning two server classes (H100 SXM and A100 SXM) with multiple tensor parallel sizes, three open-source models (Llama3.1-8B [13], Qwen2.5-14B/32B [17]), and workloads with 128–2048 concurrent requests and heterogeneous sequences that push total tokens to ninety percent of the engine KV-cache capacity (C). We cover both prefill and decode, report coefficient of determination (r^2) and relative error at the 90th and 99th percentiles, and compare against VTC [16], a token based

Model	Llama3.1-8B			Qwen2.5-14B			Qwen2.5-32B		
Parallelism	A1	H1	H4	H1	H2	H4	A2	H2	H4
VTC R^2	0.998	0.998	0.995	1.000	0.997	0.998	1.000	1.000	1.000
VTC Err _{p90}	0.08	0.09	0.04	0.03	0.14	0.08	0.02	0.03	0.02
VTC Err _{p99}	0.50	0.79	0.16	0.09	1.24	0.57	0.10	0.07	0.11
<i>LLMVisor</i> R^2	1.000	0.999	1.000	1.000	1.000	1.000	1.000	1.000	1.000
<i>LLMVisor</i> Err _{p90}	0.03	0.10	0.01	0.02	0.02	0.01	0.01	0.02	0.03
<i>LLMVisor</i> Err _{p99}	0.14	0.92	0.08	0.08	0.08	0.13	0.05	0.07	0.08

Table 1: Prefill step latency modeling accuracy of VTC and *LLMVisor*.

Model	Llama3.1-8B			Qwen2.5-14B			Qwen2.5-32B		
Parallelism	A1	H1	H2	H1	H2	H4	A2	H2	H4
VTC R^2	0.947	0.958	0.778	0.991	0.948	0.938	0.972	0.991	0.992
VTC Err _{p90}	0.23	0.24	0.42	0.12	0.24	0.27	0.15	0.12	0.11
VTC Err _{p99}	0.50	0.54	0.88	0.18	0.45	0.63	0.29	0.20	0.31
<i>LLMVisor</i> R^2	0.974	0.997	0.995	0.998	0.997	0.996	0.979	0.997	0.999
<i>LLMVisor</i> Err _{p90}	0.10	0.04	0.05	0.04	0.04	0.07	0.11	0.06	0.03
<i>LLMVisor</i> Err _{p99}	0.16	0.08	0.09	0.06	0.07	0.10	0.16	0.10	0.05

Table 2: Decode step latency modeling accuracy of VTC and *LLMVisor*.

proxy that does not model batch size or context length. Further details of the evaluation setup are provided in Sec. A.

3.2 Prefill Latency Modeling

The performance of *LLMVisor* and the VTC baseline on the prefill latency modeling task is presented in Table 1. Across all evaluated models and hardware configurations, *LLMVisor* consistently outperforms VTC in modeling prefill latency.

While VTC achieves reasonably high R^2 values (> 0.995 in most cases), it exhibits much larger relative errors, particularly at p99. Averaged across all configurations (excluding the two blue-marked outliers), VTC’s relative error is 0.05 at p90 and 0.30 at p99, while *LLMVisor* reduces these to 0.02 and 0.09, respectively. This corresponds to an average improvement of roughly $2.5\times$ at p90 and over $3.3\times$ at p99. The high R^2 values for both methods indicate that VTC’s linear token-based model can capture coarse latency trends in prefill phase. However, by taking into account the square of input tokens, which reflects the self-attention computational complexity, *LLMVisor* is far more reliable for predicting prefill latency in multi-tenant serving environments.

3.3 Decode Latency Modeling

The performance of *LLMVisor* and the VTC baseline on the decode latency modeling task is presented in Table 2. Compared to the prefill phase, decode latency is substantially more challenging to predict. This difficulty is evident in VTC’s results: while it achieves moderately strong correlation in most cases ($R^2 > 0.9$), it completely breaks down in others. For example, Llama3.1-8B on H2 yields only $R^2 = 0.778$, indicating that VTC is not even able to capture the overall latency trend. At the same time, its relative errors remain high across the board, averaging 0.21 at p90 and 0.44 at p99, and reaching as high as 0.88 at p99. These results show that a simple linear token-based model is fundamentally inadequate for decode latency, as it ignores the effects of context length and batch size.

In contrast, *LLMVisor* consistently maintains near-perfect fits, with $R^2 > 0.97$ in all cases and typically above 0.995. More importantly, *LLMVisor* reduces relative error dramatically: average p90 error drops from 0.21 (VTC) to 0.06, and average p99 error from 0.44 to 0.10, corresponding to $3.5\times$ and $4.4\times$ improvements, respectively. Even in the more challenging decode phase, where latency is influenced by batching variability and sequence divergence, *LLMVisor* reliably captures both the overall trend and high-percentile behavior.

Overall, *LLMVisor* significantly outperforms linear token-based modeling in both prefill and decode phases, with up to $3.5\times$ and $4.4\times$ improvement on p90 and p99 relative error, respectively. It consistently predicts latency accurately and reliably, which will benefit the scheduling in multi-tenant serving environments.

References

- [1] Amazon Web Services. Amazon EC2 Instance Types – Burstable Performance Instances (T3). Online, 2023.
- [2] Anthropic. Anthropic API. <https://www.anthropic.com/api>, 2025. [Online].
- [3] Paul Barham, Boris Dragovic, Keir Fraser, Steven Hand, Tim Harris, Alex Ho, Rolf Neugebauer, Ian Pratt, and Andrew Warfield. Xen and the art of virtualization. *ACM SIGOPS operating systems review*, 37(5):164–177, 2003.
- [4] Rishi Bommasani, Drew A Hudson, Ehsan Adeli, Russ Altman, Simran Arora, Sydney von Arx, Michael S Bernstein, Jeannette Bohg, Antoine Bosselut, Emma Brunskill, et al. On the opportunities and risks of foundation models. *arXiv preprint arXiv:2108.07258*, 2021.
- [5] Aaron Hurst, Adam Lerer, Adam P Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, et al. Gpt-4o system card. *arXiv preprint arXiv:2410.21276*, 2024.
- [6] Saki Imai, Rina Nakazawa, Marcelo Amaral, Sunyanan Choochotkaew, and Tatsuhiko Chiba. Predicting llm inference latency: A roofline-driven ml method. In *Annual Conference on Neural Information Processing Systems*, 2024.
- [7] Bowen Jin, Hansi Zeng, Zhenrui Yue, Jinsung Yoon, Serkan Arik, Dong Wang, Hamed Zamani, and Jiawei Han. Search-r1: Training llms to reason and leverage search engines with reinforcement learning. *arXiv preprint arXiv:2503.09516*, 2025.
- [8] Shuowei Jin, Xueshen Liu, Qingzhao Zhang, and Z Morley Mao. Compute or load kv cache? why not both? *arXiv preprint arXiv:2410.03065*, 2024.
- [9] Yunho Jin, Chun-Feng Wu, David Brooks, and Gu-Yeon Wei. $\mathcal{G}$: Increasing gpu utilization during generative inference for higher throughput. *Advances in Neural Information Processing Systems*, 36:18015–18027, 2023.
- [10] Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E. Gonzalez, Hao Zhang, and Ion Stoica. Efficient memory management for large language model serving with PagedAttention. In *Proceedings of the 29th ACM Symposium on Operating Systems Principles (SOSP)*, 2023. arXiv:2309.06180.
- [11] Paul Menage, Paul Jackson, and Christoph Lameter. Cgroups. Available on-line at: <http://www.mjmwired.net/kernel/Documentation/cgroups.txt>, 2008.
- [12] Dirk Merkel et al. Docker: lightweight linux containers for consistent development and deployment. *Linux j*, 239(2):2, 2014.
- [13] Meta. Introducing llama 3.1: The next generation of open models. <https://ai.meta.com/blog/meta-llama-3-1/>, July 2024. Accessed: 2025-03-29.
- [14] NVIDIA Corporation. NVIDIA Multi-Instance GPU (MIG) User Guide. <https://docs.nvidia.com/datacenter/tesla/mig-user-guide/>, 2025. [Online].
- [15] OpenAI. OpenAI API. <https://openai.com/api/>, 2025. [Online].
- [16] Ying Sheng, Shiyi Cao, Dacheng Li, Banghua Zhu, Zhuohan Li, Danyang Zhuo, Joseph E Gonzalez, and Ion Stoica. Fairness in serving large language models. In *18th USENIX Symposium on Operating Systems Design and Implementation (OSDI 24)*, pages 965–988, 2024.
- [17] Qwen Team. Qwen2.5: A party of foundation models, September 2024.
- [18] Zhihang Yuan, Yuzhang Shang, Yang Zhou, Zhen Dong, Zhe Zhou, Chenhao Xue, Bingzhe Wu, Zhikai Li, Qingyi Gu, Yong Jae Lee, et al. Llm inference unveiled: Survey and roofline model insights. *arXiv preprint arXiv:2402.16363*, 2024.

189 A Evaluation Setup Details

190 Our evaluation framework is built on vLLM (v0.7.3) [10], using its internal profiler to capture ground
191 truth latency for each processing step. To ensure a comprehensive assessment, we systematically vary
192 the hardware, models, and workloads.

193 **Hardware** We conduct experiments on two distinct server types with five different tensor parallelism
194 settings to cover various deployment scenarios: a. A server with 4 NVIDIA H100 SXM 80GB GPUs,
195 a 104-core vCPU, and 1800GiB of memory. Tensor parallel size 1, 2, and 4, denoted as **H1**, **H2**, **H4**.
196 b. A server with 8 NVIDIA A100 SXM 80GB GPUs, a 240-core vCPU, and 1800GiB of memory.
197 Tensor parallel size 1 and 2, denoted as **A1**, **A2**.

198 **Models** We test *LLMVisor* on three popular open-source LLMs of varying sizes to evaluate its
199 generalizability: Llama3.1-8B [13], Qwen2.5-14B [17], and Qwen2.5-32B [17].

200 **Workloads** To simulate realistic, multi-tenant serving environments, we generate workloads with
201 varying numbers of concurrent requests, from 128 to 2048. To test performance across different
202 context lengths, we increase the total number of tokens from 4096 to 90% of the engine’s maximum
203 KV cache capacity (C). These tokens are then randomly distributed among requests with high variance
204 to ensure a heterogeneous mix of request lengths. Our evaluation covers both the prefill and decode
205 phases of inference.

206 **Metrics** We assess the accuracy of resource modeling methods via two metrics: a. Coefficient
207 of Determination (r^2): Quantifies the linear correlation between the predicted and ground truth
208 latencies. A higher r^2 indicates a better fit. b. Relative Error: Measures the percentage difference
209 between predicted and actual latency, reported at the 90th (p90) and 99th (p99) percentiles to evaluate
210 performance.

211 **Baseline** We benchmark *LLMVisor* against VTC [16], a baseline model that estimates request latency
212 based on a linear model of token number. Unlike *LLMVisor*, VTC does not account for critical
213 system-level factors such as batch size and context length.