

Validation of vLLM PagedAttention Throughput and Memory-Efficiency Claims

A desk-based systems evaluation of mechanism validity, benchmark evidence, interpretation, reproducibility, and technical risk

Organization	Quant Labs
Evaluation date	July 31, 2026
Evaluation type	Desk-based technical verification
System evaluated	vLLM
Core technology	PagedAttention / KV-cache memory management
Evaluation status	Claim supported, with qualification; not independently reproduced

Research-paper scope

This paper evaluates whether the public technical claim behind PagedAttention is internally coherent, falsifiable, and supported by the disclosed benchmark evidence. It is **not** a Quant Labs benchmark run on hardware under our control. The distinction between source validation and independent experimental reproduction is maintained throughout the paper.

SUPPORTED Technical mechanism	QUALIFIED Headline multiplier	PENDING Independent reproduction	HIGH Production relevance
---	---	--	-------------------------------------

Abstract

Large-language-model (LLM) serving systems must retain key-value (KV) tensors for active sequences during autoregressive decoding, making KV-cache allocation a first-order GPU-memory constraint. vLLM’s PagedAttention architecture addresses this problem by dividing KV-cache storage into fixed-size blocks and mapping logical sequence blocks to non-contiguous physical memory. This paper performs a desk-based technical evaluation of the architecture and the throughput claims associated with its original public release and SOSP 2023 paper. The evaluation separates three questions that are frequently conflated: whether the mechanism is technically credible, whether the reported benchmarks are interpretable, and whether a headline multiplier can be generalized beyond its tested baseline. We find that the mechanism is technically sound and falsifiable, the original benchmark methodology is sufficiently disclosed for technical interpretation, and the reported results support material throughput improvement relative to the evaluated baselines. However, the widely cited “up to 24×” result applies to a basic Hugging Face Transformers serving baseline and should not be interpreted as a universal vLLM advantage. The SOSP 2023 paper reports a smaller but more meaningful approximately 2–4× improvement relative to optimized serving systems used as research baselines, depending on workload and model configuration. Quant Labs has not independently reproduced either multiplier on controlled hardware. We therefore classify the core claim as *supported with qualification* and propose a controlled reproduction protocol centered on sustainable throughput under latency service-level objectives (SLOs), rather than aggregate token throughput alone.

Keywords: vLLM; PagedAttention; large language model serving; KV cache; GPU memory; inference throughput; batching; TTFT; TPOT; technical due diligence.

1 Introduction

LLM inference is a systems problem as much as a model-compute problem. During autoregressive generation, each active sequence retains attention key and value tensors that grow with sequence length. When an inference server handles many simultaneous requests, the KV cache can consume a significant fraction of available accelerator memory. How that memory is allocated determines how many sequences can remain resident, how effectively requests can be batched, and ultimately how much useful throughput a serving system can sustain.

vLLM introduced PagedAttention as a block-based KV-cache allocation architecture [2]. Rather than requiring the KV cache for a sequence to occupy a single large contiguous region, the system divides the cache into fixed-size blocks and maps logical blocks to available physical blocks. The design is conceptually analogous to paging in virtual memory: logical sequence order is preserved even when physical storage is non-contiguous.

The resulting claim is straightforward: reducing KV-cache fragmentation and unnecessary reservation should increase effective GPU-memory utilization; higher usable memory should permit more concurrent sequences; increased concurrency should enable larger or more efficient continuous batches; and those batches should raise aggregate serving throughput.

1.1 Research Question

This evaluation asks:

Research question

Does the public evidence support the claim that PagedAttention materially improves LLM-serving throughput by improving KV-cache memory efficiency, and how should the reported performance multipliers be interpreted?

1.2 Scope and Contributions

The paper makes five contributions:

1. isolates the systems mechanism behind the claim;
2. distinguishes memory-efficiency gains from raw model-compute acceleration;
3. evaluates the interpretability of the original benchmark methodology;
4. separates baseline-specific performance multipliers from universal claims; and
5. specifies a minimum controlled reproduction protocol suitable for technical due diligence.

Critical scope limitation

Quant Labs has not independently reproduced the original benchmark. This document validates the technical mechanism and the interpretability of public evidence; it does not certify the exact reported multipliers under current hardware, software, models, or workloads.

2 Technical Background: Autoregressive Inference and the KV Cache

Autoregressive transformer inference can be divided into two operational phases: prefill and decode. The two phases stress the serving system differently.

2.1 Prefill

During prefill, the model processes the input prompt and generates KV-cache entries corresponding to the prompt tokens. This phase is comparatively compute-intensive and can exploit substantial parallelism.

2.2 Decode

During decode, output tokens are generated sequentially. At each step, the attention mechanism must reference the sequence history, so previously generated key and value tensors remain available. The KV cache therefore grows as the sequence grows.

For a serving system with many active requests, KV-cache capacity can become a dominant constraint on the number of sequences that remain resident in GPU memory. This makes memory allocation a throughput concern, not merely an implementation detail.

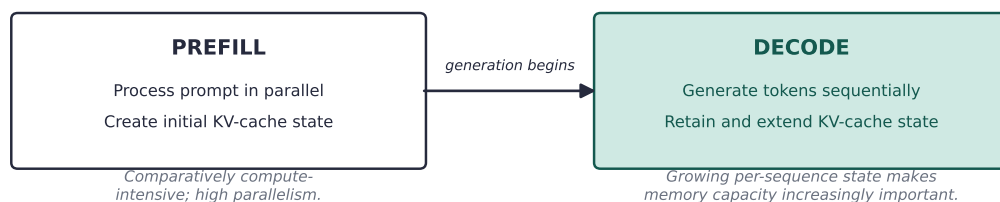


Figure 1: Operational distinction between prefill and decode. PagedAttention primarily addresses the memory-management consequences of maintaining growing KV-cache state across concurrent sequences.

3 The Memory-Allocation Problem

A straightforward serving implementation may reserve contiguous memory for each request’s KV cache. Variable sequence lengths and unknown future generation lengths create four related inefficiencies.

3.1 Unknown Output Length

At request admission, the server does not know exactly how many tokens will be generated. Reserving for a maximum sequence length wastes memory when requests terminate early. Reserving too conservatively can require reallocation or restrict concurrency.

3.2 Internal Fragmentation

A request can hold memory that it never uses when its actual sequence is shorter than the allocated capacity. The unused region is unavailable to other sequences even though it contributes no useful state.

3.3 External Fragmentation

Sufficient aggregate free memory may exist while being split into smaller non-contiguous regions. If the allocator requires a large contiguous region, fragmented capacity can prevent admission of a new sequence.

3.4 Reduced Batching Capacity

Memory stranded by inefficient allocation cannot hold additional active sequences. This lowers the number of requests that can participate in a continuous batch and therefore reduces aggregate serving throughput.

4 PagedAttention Architecture

PagedAttention changes the allocation problem from finding one sufficiently large contiguous region to finding enough available fixed-size KV-cache blocks. A sequence sees a logically contiguous cache, while its physical blocks may be distributed across GPU memory.

Traditional contiguous reservation



PagedAttention block allocation



Figure 2: Conceptual comparison of contiguous reservation and PagedAttention-style block allocation. Contiguous reservation can strand unused capacity and fragmented free regions; block-based allocation permits logical sequences to occupy non-contiguous physical KV-cache blocks. The figure is schematic, not a measured memory trace.

A conceptual block table for a single sequence can be written as:

Logical block	Physical block	Implication
Sequence A – Block 0	Physical Block 14	First logical chunk may occupy any free block.
Sequence A – Block 1	Physical Block 37	No physical adjacency with Block 0 is required.
Sequence A – Block 2	Physical Block 8	Logical order is recovered through mapping.
Sequence A – Block 3	Physical Block 52	Growth consumes additional reusable blocks.

4.1 Expected Performance Mechanism

PagedAttention does not inherently reduce the transformer’s floating-point operations per token. Its principal benefit is better serving-system utilization. The expected mechanism is shown in Figure 3.



Figure 3: Expected performance mechanism. The throughput gain is a systems-utilization effect, not a claim that model computation itself becomes 24× faster.

Technically precise interpretation

A defensible statement is that PagedAttention can increase serving throughput by using available GPU memory more efficiently and allowing the scheduler to maintain higher request concurrency. The statement “PagedAttention makes model computation 24× faster” would be misleading.

5 Evaluation Methodology

This evaluation uses source and architecture validation rather than independent benchmark reproduction. The methodology has four components: primary-source review, mechanism analysis, benchmark-interpretability assessment, and corroboration assessment.

5.1 Primary Evidence Reviewed

The source report identifies two principal evidence classes:

1. the 2023 vLLM launch material describing “Easy, Fast, and Cheap LLM Serving with PagedAttention” [1]; and
2. Kwon et al., “Efficient Memory Management for Large Language Model Serving with PagedAttention,” published at SOSP 2023 [2].

The research paper describes the serving architecture, KV-cache allocation strategy, comparison baselines, workload assumptions, model configurations, hardware environment, scheduling behavior, and throughput measurements. This disclosure makes the experiments technically interpretable even though they have not been rerun by Quant Labs.

5.2 Evaluation Criteria

We evaluate the claim against the following questions:

1. Is the mechanism causally plausible?
2. Are the measurable consequences of the mechanism falsifiable?
3. Are the benchmark baselines and configurations described sufficiently to interpret the reported numbers?
4. Does the reported improvement match the expected direction of the architectural change?
5. Is the headline multiplier being used only within the baseline and workload scope that produced it?

5.3 What Desk-Based Verification Can and Cannot Establish

Desk-based verification can establish whether a mechanism is coherent, whether a claim is testable, and whether the disclosed evidence supports a bounded conclusion. It cannot establish current performance parity, exact replication, or superiority across workloads not tested in the source material.

6 Results

6.1 Mechanism Validity

The mechanism is technically credible. KV-cache fragmentation is a legitimate serving constraint; block-based allocation directly targets that constraint; higher memory efficiency can increase the number of concurrently resident sequences; and increased concurrency can improve aggregate throughput through more effective batching.

6.2 Benchmark Transparency

The original research provides enough methodological detail to make the performance results interpretable. This is important because inference results are highly sensitive to hardware, model,

workload, scheduler, sequence-length distribution, batching policy, and latency objectives. A result that omitted these factors would carry materially less evidentiary weight.

6.3 Reported Performance Is Baseline-Specific

Two performance claims should be kept separate. The vLLM launch material reported up to approximately $24\times$ higher throughput than a basic Hugging Face Transformers serving implementation [1]. The SOSP 2023 paper reported approximately $2\text{--}4\times$ higher throughput than optimized research baselines such as Orca and FasterTransformer, depending on workload and model configuration [2].

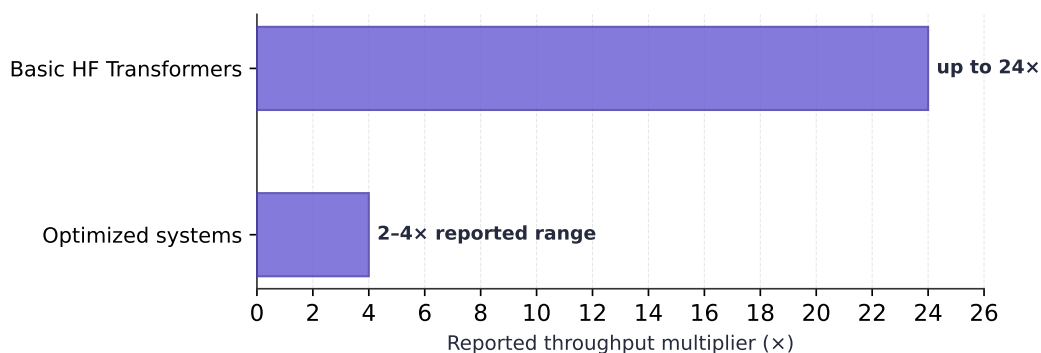


Figure 4: Baseline-specific interpretation of reported throughput improvements. The optimized-system bar is drawn at the upper bound of the reported $2\text{--}4\times$ range for visual comparison; it is not a claim of a constant $4\times$ result.

Interpretation rule

The phrase “vLLM is $24\times$ faster” removes the baseline and therefore overstates what was demonstrated. A technically defensible formulation is: under the configurations reported in its original evaluation, vLLM achieved up to approximately $24\times$ the throughput of a basic Hugging Face Transformers baseline and approximately $2\text{--}4\times$ the throughput of several optimized serving systems used as research baselines.

Comparison	Reported result	Interpretation
vLLM vs. basic Hugging Face Transformers serving	Up to $24\times$	Large improvement against a relatively basic baseline; not a universal characteristic.
vLLM vs. contemporary optimized serving systems	Approximately $2\text{--}4\times$ in reported workloads	More meaningful systems-level comparison within the tested configurations.
vLLM vs. every modern inference engine	Not established	Requires fresh, controlled benchmarking against current alternatives.

Table 1. Interpretation of the reported performance results.

6.4 Why the Result Is Technically Plausible

PagedAttention is most valuable when GPU memory constrains concurrency. If one server can keep only a limited number of active sequences resident because substantial KV capacity is stranded, while a more efficient allocator can keep materially more sequences resident on the same GPU, the second system can sustain higher aggregate throughput even when the model operations required for each token are essentially unchanged. This is a serving-system utilization improvement rather than a direct reduction in per-token matrix-multiplication cost.

7 Benchmark Variables That Must Be Controlled

Inference benchmarks are configuration-sensitive. At minimum, a credible comparison should control the factors below.

Category	Variables
Hardware	GPU model and count; GPU memory capacity; interconnect architecture; CPU configuration; host memory; PCIe / NVLink configuration.
Model	Model architecture; parameter count; numerical precision; quantization method; tensor-parallel configuration; pipeline-parallel configuration.
Workload	Prompt-length distribution; output-length distribution; request arrival rate; concurrent-user count; sampling configuration; number of returned sequences.
Serving configuration	Batch scheduler; maximum batch-token count; KV-cache block size; memory-utilization limit; prefix caching; speculative decoding; CUDA-graph usage; attention kernel; quantization back-end.

Table 2. Minimum benchmark-control variables identified in the source evaluation.

Without controlling these variables, two throughput numbers may not be directly comparable.

8 Throughput Alone Is Not a Sufficient Production Metric

Aggregate throughput does not fully characterize serving quality. A production evaluation should measure capacity and latency simultaneously.

A system can report excellent average throughput while delivering poor P95 or P99 latency under saturation. For production serving, that tradeoff is often more important than the peak token-rate number.

Metric	Typical unit / definition	Why it matters
Request throughput	requests / second	Complete-request serving capacity.
Token throughput	output tokens / second, or total processed tokens / second	Aggregate serving capacity; methodology must define token counting.
Time to first token (TTFT)	request arrival → first generated token	Critical for interactive responsiveness.
Time per output token (TPOT)	decode latency after generation begins	Characterizes generation pace.
Inter-token latency (ITL)	delay between successive streamed tokens	Directly affects streaming smoothness.
End-to-end latency	request arrival → completion	User-perceived total request time.
Tail latency	P95 / P99 variants of latency measures	Captures degradation under saturation that averages can hide.

Table 3. Core serving metrics.

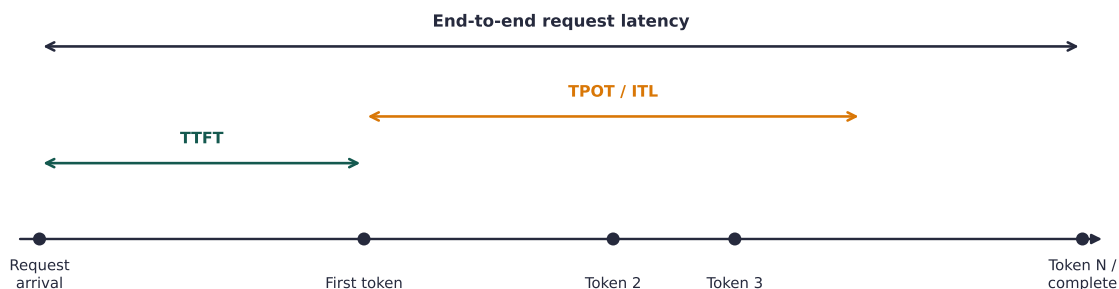


Figure 5: Latency timeline. High token throughput can coexist with poor TTFT or tail latency, so sustainable production performance must be evaluated under explicit latency SLOs.

9 Falsifiability and Corroboration

The central PagedAttention hypothesis is unusually falsifiable for an infrastructure claim. If KV-cache allocation wastes less GPU memory, then, all else equal, a server should support greater concurrency before exhausting memory. The mechanism can be evaluated by measuring GPU-memory utilization, KV-cache utilization, maximum active sequences, scheduler queue depth, batch size, request throughput, token throughput, TTFT, and TPOT under identical workloads.

The source evaluation also treats continued real-world adoption of vLLM as corroborating evidence that the platform provides practical technical value. Adoption, however, is not benchmark evidence. It can be driven by API compatibility, model support, ecosystem integrations, developer experience, community size, hardware compatibility, and operational tooling. Accordingly, adoption increases confidence in practical relevance but does not independently validate a specific throughput multiplier.

10 Evaluation Verdict

The core PagedAttention thesis survives desk-based technical review.

Dimension	Assessment
Technical mechanism	Strong
Claim falsifiability	Strong
Benchmark transparency	Strong
Evidence for material throughput improvement	Strong
Evidence for universal 24× improvement	Weak / not claimed
Independent Quant Labs reproduction	Not completed
Evidence of enduring architectural moat	Uncertain
Production relevance	High

Table 4. Final assessment across evaluation dimensions.

10.1 Supported

The review supports the following statements:

- KV-cache fragmentation is a legitimate LLM-serving constraint.
- Block-based KV-cache allocation is technically well motivated.
- Increased memory efficiency can increase concurrency.
- Increased concurrency can materially increase serving throughput.
- The original research discloses enough methodology to make its experiments technically interpretable.
- Significant performance gains were reported relative to the baselines evaluated in the original work.

10.2 Requires Qualification

The 24× figure is baseline-specific. It should not be used as shorthand for a universal 24× advantage over competing inference systems. The more useful engineering conclusion is that PagedAttention materially improved the memory efficiency and throughput of LLM serving relative to the systems evaluated at the time, with the magnitude of improvement depending heavily on workload and baseline.

10.3 Not Yet Verified

Quant Labs has not independently verified:

- the exact $24\times$ multiplier;
- the exact $2\text{--}4\times$ multiplier under current software versions;
- current superiority against modern inference engines;
- performance across every model architecture, latency objective, or GPU architecture; or
- cost efficiency under a production cloud deployment.

11 Recommended Independent Reproduction Protocol

If vLLM were the core technology of a portfolio-stage company, reported throughput figures should be replaced by a controlled benchmark. The minimum protocol below converts the source evaluation into an experimental design.

11.1 Systems Under Test

- **System A:** vLLM.
- **System B:** strongest technically relevant competing serving engine.
- **System C:** reference implementation or conventional baseline.

11.2 Fixed Configuration

The systems should use the same GPU hardware, model weights, precision, quantization, tensor parallelism, prompt distribution, output distribution, request-arrival pattern, and latency requirements.

11.3 Load Sweep

Request rates should be increased through approximately 25%, 50%, 75%, 100%, 125%, and 150% of estimated saturation throughput. At each load point, the benchmark should record both capacity and latency.



Figure 6: Recommended load sweep. The purpose is to observe not only peak throughput but the point at which latency, queue depth, or error rate becomes unacceptable.

11.4 Metrics to Record

At each point, record requests/s, tokens/s, median/P95/P99 TTFT, median/P95/P99 TPOT, GPU utilization, GPU-memory utilization, KV-cache utilization, maximum concurrent sequences, error rate, and scheduler queue depth.

11.5 Decision Rule

The critical question is not simply which engine produces the most tokens per second. A stronger production criterion is:

Decision rule

Which engine delivers the highest sustainable throughput while remaining inside the target latency SLO?

This criterion is materially harder to optimize and more representative of production inference economics.

12 Remaining Technical Risks

12.1 Competitive Convergence

PagedAttention's underlying insight is increasingly part of the broader inference-engine design space. A successful optimization does not necessarily create a durable moat once competing systems adopt equivalent memory-management concepts. The historical technical importance of an architecture and its forward-looking differentiation are therefore separate questions.

12.2 Bottleneck Migration

Better KV-cache utilization can move the system bottleneck elsewhere. Depending on workload, performance may become constrained by memory bandwidth, attention kernels, matrix multiplication, inter-GPU communication, scheduler overhead, prefill compute, decode compute, or network I/O. An optimization can remain technically effective while contributing a smaller percentage of total system performance over time.

12.3 Workload Dependence

Relative performance is unlikely to be constant across short versus long prompts, short versus long generations, low versus high concurrency, latency-sensitive serving, and offline batch inference. The architecture should be evaluated under the workload a customer actually intends to run.

13 Discussion

13.1 What This Evaluation Establishes

The strongest characteristic of the PagedAttention claim is its measurable causal chain: less KV-cache waste should increase memory utilization; greater usable memory should increase concurrency; higher concurrency should improve sustainable batching; and more effective batching should improve throughput. Each stage can be instrumented. That makes the claim suitable for meaningful technical due diligence rather than relying solely on a headline benchmark.

13.2 What This Evaluation Does Not Establish

This evaluation does not establish that vLLM is currently the best inference engine for every production workload. It also says nothing about model accuracy, hallucination rate, reasoning capability, training quality, safety, fine-tuning quality, or application-level correctness. PagedAttention operates at the inference-serving systems layer; its purpose is to improve resource utilization and serving efficiency, not model intelligence.

13.3 Implications for Technical Due Diligence

If a company claimed that a proprietary inference architecture delivered materially higher throughput, the correct diligence standard would be a controlled benchmark against the strongest relevant baseline, not only a basic reference implementation. Hardware, model, precision, workload, concurrency, sequence lengths, scheduler configuration, and latency SLOs should be controlled. The company should demonstrate improvement in both system efficiency and production-relevant performance.

The thesis would weaken if competing engines achieved equivalent memory utilization, PagedAttention ceased to provide a material advantage, other bottlenecks dominated inference cost, new model architectures materially changed KV-cache requirements, or specialized hardware reduced the economic importance of software-level cache allocation.

14 Conclusion

PagedAttention represents a credible and important systems-level optimization for LLM inference. The public evidence supports the conclusion that improved KV-cache memory management can materially increase serving throughput by increasing memory efficiency and request concurrency. The architecture is technically coherent, its causal mechanism is measurable, and the original benchmark methodology is sufficiently disclosed to interpret the reported results.

The frequently repeated $24\times$ number, however, should be treated as a benchmark-specific result against a basic Hugging Face Transformers serving baseline, not as a universal performance characteristic of vLLM. The approximately $2\text{--}4\times$ gains reported against optimized serving systems in the original research provide a more meaningful systems-level comparison, but they are still workload- and configuration-dependent.

Accordingly, the evaluation verdict is **supported with qualification**. The architecture passes desk-based technical validation; the exact benchmark multipliers remain subject to independent reproduction. For production or investment diligence, the next step should be a controlled benchmark that optimizes for sustainable throughput within explicit latency SLOs and compares vLLM against the strongest current alternative under identical conditions.

A Minimum Reproduction Matrix

Dimension	Minimum requirement
Systems	vLLM; strongest relevant competitor; conventional/reference baseline.
Hardware	Identical GPU model/count/memory, CPU, host memory, and interconnect.
Model	Same weights, architecture, precision, quantization, tensor/pipeline parallel settings.
Workload	Controlled prompt/output distributions, arrival process, concurrency, sampling, returned sequences.
Serving	Controlled scheduler settings, batch-token cap, memory target, KV block size, cache options, kernels, CUDA graphs.
Load	25%, 50%, 75%, 100%, 125%, and 150% of estimated saturation.
Capacity metrics	Requests/s; output or total tokens/s with methodology explicitly defined.
Latency metrics	Median/P95/P99 TTFT; median/P95/P99 TPOT; ITL where streaming is relevant; end-to-end latency.
Resource metrics	GPU utilization; GPU-memory utilization; KV-cache utilization; max concurrent sequences; queue depth.
Reliability metrics	Error rate and request failures under load.
Decision rule	Highest sustainable throughput while remaining inside the target latency SLO.

Table 5. Suggested minimum experimental matrix for an independent reproduction.

B Evaluation Summary for Technical Reviewers

One-paragraph reviewer summary

PagedAttention’s technical thesis is credible because it targets a concrete source of GPU-memory inefficiency in autoregressive serving and produces a measurable sequence of expected effects: lower KV-cache waste, higher effective memory use, greater sequence concurrency, and higher sustainable throughput. Public evidence supports material improvement relative to the tested baselines, but the multiplier depends on workload and comparison system. The “up to 24×” launch result should not be generalized beyond its basic Hugging Face Transformers baseline. Quant Labs has not rerun the benchmark, so the present result is best described as source-supported and architecture-validated rather than independently reproduced.

References

1. vLLM Project. “vLLM: Easy, Fast, and Cheap LLM Serving with PagedAttention.” Project launch material, June 20, 2023.

2. W. Kwon, Z. Li, S. Zhuang, et al. “Efficient Memory Management for Large Language Model Serving with PagedAttention.” *Proceedings of the 29th ACM Symposium on Operating Systems Principles (SOSP ’23)*, 2023.
3. vLLM Project. Benchmark suites and performance-dashboard documentation.
4. vLLM Project. Benchmark scripts and reference implementations.

Quant Labs. Technical and financial due diligence in AI, robotics, and quantum computing. This document is for general informational purposes only and does not constitute investment advice. Source report: Quant Labs Technical Evaluation #001.