

TreeWY: Speculative Verification for Gated DeltaNet Hybrids

Sneha Murthy Ghantasala

Thomson Reuters

sneha.ghantasala@thomsonreuters.com

Abstract

Modern open models are *hybrids*: most layers are linear-attention (Gated DeltaNet, GDN) layers carrying a small fixed-size recurrent state instead of a growing key-value (KV) cache. This makes ordinary decoding memory-efficient, but hurts speculative decoding. To verify a batch of draft tokens and then roll back the rejected ones, today’s systems snapshot the full recurrent state at *every* draft position for GDN layers, and those snapshots cannot be shared across branches of a draft tree, so a wide, high-acceptance tree becomes memory-infeasible. We remove the snapshots. Using a *tree-structured WY transform* of the gated delta rule, we compute every draft node’s output with a single triangular solve and reconstruct *only* the one accepted state on commit, storing a small pseudo-value matrix instead of per-node states; the derivation depends only on the gated delta rule, not on any other architectural detail. In serving benchmarks on two scales of one hybrid model family (Qwen3.5 35B and 397B) this cuts speculative recurrent-state memory and KV-cache pressure at identical acceptance length, turning the freed HBM into higher throughput and much lower time-to-first-token (TTFT) wherever memory binds, and costing a few percent where it does not. For tree width the same memory buys affordability: a wider, higher-acceptance draft becomes possible, though not yet a throughput win.

1 Background

Speculative decoding. Autoregressive decoding emits one token per forward pass, so it is memory-bandwidth bound and the GPU sits mostly idle. Speculative decoding [1, 2] fills that idle compute: a cheap drafter (a small model, or a lightweight “MTP” head shipped with the target) proposes k candidate tokens; the target verifies all k in one pass and accepts the longest prefix it would have generated itself, leaving the output distribution unchanged but emitting several tokens per pass. A *tree* drafter proposes multiple alternatives per position [3].

Hybrid models and GDN layers. A *hybrid* model mixes softmax-attention with linear-attention layers (for e.g., Gated DeltaNet, GDN). A softmax layer keeps a KV cache: it stores a key and value vector for *every* previous token, so its memory grows linearly with context. A *Gated DeltaNet* (GDN) layer [4] instead maintains one recurrent *state* matrix $S \in \mathbb{R}^{d_v \times d_k}$ per head summarizing the whole prefix in fixed size. For token t with key k_t , value v_t , query q_t , scalar decay gate $\alpha_t \in (0, 1)$ and write strength $\beta_t \in (0, 1)$, the state evolves by the **gated delta rule**

$$S_t = \alpha_t S_{t-1} (I - \beta_t k_t k_t^\top) + \beta_t v_t k_t^\top, \quad o_t = S_t q_t. \quad (1)$$

The transition $T_t = \alpha_t (I - \beta_t k_t k_t^\top)$ is a *scalar decay* times a *rank-1 correction*: the $\beta_t k_t k_t^\top$ term “erases” what the state already predicts for k_t before writing the new value (the “delta” rule). Crucially, S is a *lossy summary* of the whole prefix. It cannot be truncated or partially rolled back the way a KV cache can. That single fact is the source of the speculative-decoding problem.

2 The memory problem

We make the contrast concrete on the two Qwen3.5 [5] hybrids we evaluate, both 3:1 GDN-to-softmax with $d_k = d_v = 128$ and softmax $d_{\text{head}} = 256$, 2 KV heads: *35B-A3B* (30 GDN, 32 value heads; 10 softmax) and *397B-A17B* (45 GDN, 64 value heads; 15 softmax).

Normal decoding. Summed over the softmax layers, the KV cache costs $2(K,V) \times n_{\text{softmax}} \times n_{\text{KV-head}} \times d_{\text{head}} \times 2$ B (bf16) per token: 20 KiB/token at 35B and 30 KiB at 397B, and it *grows with context* (0.625 and 0.94 GiB/seq at a 32K context). The GDN layers instead hold one fixed state, $n_{\text{GDN}} \times n_{\text{v-head}} \times d_v \times d_k \times 2$ B (bf16) per sequence: 30 and 90 MiB/seq, forever, which is why hybrids are attractive.

Speculative verification: where GDN blows up. Speculation inverts the picture. To verify k draft tokens (or a tree of N nodes):

- **Softmax layers stay cheap.** Draft KV entries are appended, and rejected ones are dropped by moving a pointer, i.e., rollback is free. A tree even *shares* its common-prefix KV across branches. The marginal cost is a handful of KiB.
- **GDN layers explode during verification.** Verification is one parallel forward pass over all k tokens at once. In that single pass the GDN recurrent state advances all the way to the end (past token k) before any acceptance is checked. And so, we need to rollback to the state at the accepted node. The standard fix, which we call *full-state snapshotting*, keeps the committed state plus one full snapshot per draft position: $k+1$ blocks per sequence per GDN layer for a chain of k draft tokens (a tree of N nodes needs $N+1$, one per node). For $k = 3$ that is $4\times$ the committed state, 120 MiB/seq (35B) and 360 MiB/seq (397B). Moreover, the snapshots are *unshareable across branches*. Snapshot memory scales with draft tokens, caps concurrency and rules out wide trees. This is the bottleneck we remove.

Existing approaches and trade-offs.

- **(i) Full-state snapshotting.** The current default in *both* vLLM [6] (`storeall` in our figures) and SGLang [7]: one recurrent-state snapshot per draft position, rolled back by indexing to the accepted one. Rollback is free; memory scales with draft size and snapshots cannot be shared across tree branches.
- **(ii) Deferred materialization with a rank-1 cache** (ReplaySSM [8], shipping concurrently July 2026 in a vLLM RFC and TensorRT-LLM) keeps a checkpoint plus a short cached history and defers materializing the state until a periodic flush, so state is written far less often than every draft step. Inside the verify window it solves the *same* triangular system we do (Section 3), so the distinction is *when* state is materialized, not how it is computed. It covers Mamba2 and GDN, and the implementation linked in the RFC is chain-only (not a tree). It is the closest prior point in vLLM; we benchmark against it in Appendix D. ReplaySSM has also been incorporated in SGLang [9] as an opt-in feature.
- **(iii) Transition-matrix tree recompute** (STree [10, 11]) verifies a tree, but only for *Mamba2*: its transition is a bare *scalar* decay, so its speedup rests on turning a product of gates into a cheap cumulative sum. This does *not* exist for GDN’s non-commuting matrix transition (scalar decay $\times$ rank-1 correction).
- **(iv) Tree-structured closed form in SGLang** (Bole [12], concurrent, Aug. 2026) is the most direct overlap with this paper: it also rewrites the linear-attention recurrence into a tree-structured closed form so a whole draft tree can be verified in parallel rather than branch by branch, reporting $82\text{--}99\times$ transient-memory reduction and $3.4\text{--}7.7\times$ faster tree verification from its kernel. It targets SGLang rather than vLLM and, by its own description, is framed at the level of hybrid-attention recurrences generally. We have not run a head-to-head, since Bole’s paper does not link to a code implementation; it is the closest concurrent work and the comparison we most want to make next.

Our method is the tree-structured WY/UT transform of the gated delta rule (Section 3), implemented in vLLM (Section 4): one triangular solve verifies every draft node. It generalizes DeltaNet’s [13] chain WY transform to a tree strict-ancestor form reused for verification and rollback. A single small pseudo-value matrix written only on commit stands in for the per-node snapshots. (iv) is architecturally closest to us and was arrived at independently on a different serving stack.

3 Method: TreeWY

Key rewrite. Expanding (1) shows the delta rule is *decayed additive attention with a corrected value*: $S_t = \alpha_t S_{t-1} + \tilde{v}_t k_t^\top$ with *pseudo-value* $\tilde{v}_t = \beta_t(v_t - \alpha_t S_{t-1} k_t)$ (raw value minus what the state already predicts). Given the $\tilde{v}_t$, both the state and every output are decay-weighted sums over the preceding tokens, so obtaining the $\tilde{v}_t$ without walking the recurrence removes the need for intermediate states.

Verifying a draft. A draft can be a chain or a tree: a chain proposes one candidate per position, a tree lets a node fan out into several. Either way, we lay out its N nodes in DFS pre-order from the committed state S_0 , so every ancestor of node t has a smaller index ($i \prec t$; on a chain this is just $i < t$, and $N=k$). Chaining the per-token rewrite along every root-to-node path turns verification into one linear system for all the pseudo-values at once, $\tilde{V} \in \mathbb{R}^{N \times d_v}$:

$$(I + \text{diag}(\beta) G) \tilde{V} = R, \quad G[t, i] = \frac{g_t}{g_i} (k_t^\top k_i) [i \prec t], \quad R[t] = \beta_t v_t - \beta_t g_t (S_0 k_t), \quad (2)$$

with g_t the cumulative decay from S_0 to t . Parents precede children, so G is strictly lower-triangular, and the whole draft solves in one forward substitution: no recurrence, no per-node state, and no change to the system as the tree widens. This is DeltaNet’s [13] WY/UT transform applied to the whole verify window at once; every node’s output reads out from that same solve.

Reconstruct on commit. Once verification accepts node a , we rebuild the continuation state directly from $\tilde{V}$, summing over a ’s ancestors:

$$S_a = g_a S_0 + \sum_{i \preceq a} \frac{g_a}{g_i} \tilde{v}_i k_i^\top, \quad (3)$$

which becomes the next round’s S_0 . We store only $\tilde{V}$, $O(Nd_v)$, instead of one full state per node: a $128 \times$ smaller object per head at $d_k=d_v=128$, which is what turns snapshotting’s $N+1$ state blocks per sequence per layer into one, for a chain or a tree alike. The scalar-decay ancestor mask is STree’s [10]; extending it to carry GDN’s rank-1 delta correction, so it covers GDN’s non-commuting transition rather than only Mamba2’s scalar one, is the piece we add.

4 Implementation

We implement TreeWY for vLLM, as a fork of its main branch (not yet upstreamed as of this writing) through two SpeculativeConfig options: `mamba_state_commit="reconstruct"` turns on our commit strategy in place of the default `"store_all"`, and `draft_tree_widths` turns a chain into a tree by giving the branching factor per level. All evaluations here run with prefix caching disabled. A chain (or an all-ones tree, which is really just a chain) verifies and commits in one fused, CUDA-graph-capturable Triton kernel. A real tree ($w>1$) needs a non-causal ancestor mask that cannot be replayed from a CUDA graph, so vLLM drops the whole model to piecewise capture, which also evicts the GDN mixer from graphs, making the cost far larger than the mask itself. Trees must also be scheduled *atomically*: a DFS prefix of a tree is a different topology, so a request whose tree does not fit the per-step token budget skips speculation rather than being truncated. All results use greedy drafting and verification. We check correctness two ways: the closed form matches the per-node recurrence to $\sim 10^{-15}$ (fp64) and $\sim 10^{-7}$ (fp32), and the production kernel matches that reference within bf16 tolerance. Token streams are therefore not bit-identical to the baseline; we gate correctness against a shared no-speculation reference and compare the *acceptance length* directly between `treewy` and `storeall`.

5 Evaluation

Setup. We serve Qwen3.5-35B-A3B (tensor-parallel degree 1, TP1) and Qwen3.5-397B-A17B (TP8) in vLLM on B200 GPUs (178 GiB HBM/device) with a depth-3 MTP draft chain, comparing `treewy` (ours) against snapshotting (vLLM’s default, `storeall`). We sweep GPU memory utilization (`gmu`) $\in \{0.6, 0.75, 0.9\}$ over six workloads (ShareGPT [14], spec-bench [15], BurstGPT [16], and synthetic balanced-chat, generation-heavy, summarize-heavy).¹ Five of the six sweep fixed

¹Licenses: Qwen3.5 [5] and vLLM [6] are Apache-2.0, as is spec-bench; BurstGPT is CC-BY-4.0; the ShareGPT mirror we use is tagged Apache-2.0 on Hugging Face. The synthetic workloads use vLLM’s own `random` benchmark generator.

Table 1: Concurrency sweep at each model’s tightest measured budget (35B: gmu 0.6; 397B: gmu 0.75), geomean over 5 workloads, > 1 favours TreeWY. *tput* is output throughput; *KV red.* is peak KV-cache usage reduction; *TTFT* is p99 TTFT; *TPOT* is mean time-per-output-token (TPOT). *b* is the ratio of peak *admitted* requests, the mechanism behind the other columns past the knee (remaining budgets: Table 3).

conc.	35B-A3B (TP1, gmu 0.6)					397B-A17B (TP8, gmu 0.75)				
	tput	KV red.	TTFT	TPOT	<i>b</i>	tput	KV red.	TTFT	TPOT	<i>b</i>
1	0.94	2.43	1.02	0.94	1.00	0.93	2.00	0.95	0.94	1.00
8	1.02	2.40	1.03	1.02	1.00	0.97	2.19	1.08	0.96	1.00
32	1.01	2.45	1.02	1.00	1.00	1.01	2.19	0.99	1.01	1.00
64	1.00	2.44	1.01	0.99	1.01	1.01	2.19	1.00	1.02	1.00
128	1.20	1.64	5.62	0.83	1.45	1.02	2.09	1.07	0.99	1.04
256	1.40	1.04	3.97	0.60	2.50	1.15	1.61	3.35	0.90	1.24

Table 2: Tree width vs. acceptance length (Qwen3.5-35B-A3B, TP1, gmu 0.9, spec_bench, greedy, no prefix caching; more information in Appendix B).

shape	<i>N</i>	acc.
storeall chain	3	3.24
reconstruct	3	3.23
treewy tree		
(1, 1, 1)	3	3.23
(2, 1, 1)	6	3.30
(2, 2, 1)	10	3.38
(2, 2, 2)	14	3.38
(3, 2, 2)	21	3.41
(3, 3, 2)	30	3.45
(3, 3, 3)	39	3.58

max-concurrency $\in \{1, 8, 32, 64, 128, 256\}$; BurstGPT instead sweeps its own Poisson arrival-rate $\in \{4, 8, 16, 32, 64\}$ req/s, a different independent variable.

Correctness. A *matched point* is one (dataset, concurrency, gmu) combination run under both **treewy** and **storeall**: 5 workloads $\times$ 6 concurrencies + BurstGPT $\times$ 5 arrival rates = 35 per gmu, times 3 gmu at 35B and 2 at 397B gives 175 total (105 at 35B, 70 at 397B), the same partition Fig. 3 splits by memory pressure. Across all of them, acceptance length is essentially identical to the baseline (mean $|\Delta| = 0.039$, max 0.33), matching to within 0.01 at every depth on real prompts.

Memory and throughput. TreeWY’s peak KV usage is $2\text{--}3\times$ lower at the same load (one state block per layer vs. $k+1$), so it preempts far less under pressure (1365 vs. 2531 requests; just 51 of treewy’s are at 397B, vs. 909 for storeall). That headroom drives Table 1 and Fig. 2: TreeWY wins outright where memory binds (up to $1.49\times$ throughput, $\sim 40\times$ lower p99 TTFT), and trails by a few percent on throughput ($0.97\text{--}0.99\times$) where it doesn’t, still at $2\text{--}3\times$ less KV. That per-token dip is the price of admitted concurrency, not a slower kernel (Appendix C).

The knee migrates predictably with the budget: 128 concurrency at gmu=0.6, 256 at 0.75, never at 0.9 (Appendix A); and wherever $b=1.00$ the freed memory is real but *unspent*: KV reduction still holds its full $2.2\text{--}2.5\times$ while every other column sits within 3% of parity. Fig. 3 decomposes every metric by memory pressure.

We separately benchmark against ReplaySSM on the identical 35B sweep (Appendix D).

Tree width. Table 2 sweeps tree width w at depth 3 against a store-all chain baseline (more shapes and block costs in Appendix B). TreeWY holds one state block per layer regardless of w (the same freed-memory mechanism as Table 1, now spent on width instead of concurrency), and acceptance length keeps rising with w even as the per-position rate falls (Appendix B). Width is therefore *affordable* under reconstruction where snapshotting can’t afford it at all: a store-all baseline’s per-request cost grows with N while TreeWY’s stays flat at one block. It is not yet a throughput win, though: a wider tree pushes $N+1$ tokens per step through the target and runs piecewise rather than fused (Section 4), so we report trees as enabled and correct, not as a speedup.

6 Conclusion

Treating the gated delta rule as a tree-structured WY transform lets us verify a GDN draft tree with one triangular solve and reconstruct only the accepted state on commit, trading $O(N)$ state snapshots for one small pseudo-value matrix, on chains and trees alike. On chains this is already a fused, graph-capturable kernel that turns freed HBM into higher throughput and much lower TTFT wherever memory binds, at a residual per-step cost that is an implementation artifact rather than a property of the closed form. On trees the same mechanism holds the stored state flat at one block regardless of width (Table 4), which makes a wider, higher-acceptance draft affordable, though not yet a throughput win, since acceptance is capped by draft depth and the non-capturable tree verify kernel costs more than the extra acceptance returns. Fusing the tree path into a single graph-capturable kernel and testing the deferred-state-write levers that ReplaySSM and Bole both point to are the next immediate engineering steps, alongside extending to a second model family.

References

- [1] Y. Leviathan, M. Kalman, Y. Matias. Fast Inference from Transformers via Speculative Decoding. ICML, 2023.
- [2] C. Chen et al. Accelerating Large Language Model Decoding with Speculative Sampling. arXiv:2302.01318, 2023.
- [3] T. Cai et al. Medusa: Simple LLM Inference Acceleration Framework with Multiple Decoding Heads. arXiv:2401.10774, 2024.
- [4] S. Yang, J. Kautz, A. Hatamizadeh. Gated Delta Networks: Improving Mamba2 with Delta Rule. ICLR, 2025.
- [5] Qwen Team. Qwen3.5: Towards Native Multimodal Agents. Blog, Feb. 2026. <https://qwen.ai/blog?id=qwen3.5>; model cards for the two checkpoints evaluated here at <https://huggingface.co/Qwen/Qwen3.5-35B-A3B> and <https://huggingface.co/Qwen/Qwen3.5-397B-A17B>.
- [6] W. Kwon et al. Efficient Memory Management for Large Language Model Serving with PagedAttention (vLLM). SOSP, 2023.
- [7] L. Zheng et al. SGLang: Efficient Execution of Structured Language Model Programs. NeurIPS, 2024. (Hybrid state pool as of v0.5.16, 2026.)
- [8] Dao AI Lab and NVIDIA. ReplaySSM: Cache SSM Inputs, Not State. Blog, 2026; vLLM RFC #47572, PR #47576; TensorRT-LLM PR #14203.
- [9] Y. Luo. [RFC] Porting ReplaySSM to SGLang: Faster Decode and Speculative Decoding for Hybrid (GDN/KDA) Models. sgl-project/sglang Issue #28511, 2026. <https://github.com/sgl-project/sglang/issues/28511>.
- [10] Y. Wu et al. STree: Speculative Tree Decoding for Hybrid State-Space Models. arXiv:2505.14969, 2025.
- [11] Y. Wu et al. Snakes and Ladders: Accelerating State Space Model Inference with Speculative Decoding. NeurIPS ENLSP Workshop, 2024.
- [12] L. Wang et al. Bole: Efficient Tree Speculation for Hybrid-Attention Language Models. arXiv:2608.01651, 2026.
- [13] S. Yang et al. Parallelizing Linear Transformers with the Delta Rule over Sequence Length. arXiv:2406.06484, 2024.
- [14] ShareGPT (user-shared ChatGPT conversations); we use the anon8231489123 mirror. Hugging Face. https://huggingface.co/datasets/anon8231489123/ShareGPT_Vicuna_unfiltered.
- [15] H. Xia et al. Unlocking Efficiency in Large Language Model Inference: A Comprehensive Survey of Speculative Decoding (introduces Spec-Bench). ACL Findings, 2024.
- [16] Y. Wang et al. BurstGPT: A Real-world Workload Dataset to Optimize LLM Serving Systems. arXiv:2401.17644, 2024.

A Remaining memory budgets

Table 3: The higher memory budgets, same axes and orientation as Table 1. At 397B the knee appears only at 0.75/256. Wherever b reads 1.00 the freed memory is real but unspent, as discussed in Section 5.

conc.	tput	KV red.	TTFT	TPOT	b
<i>Qwen3.5-35B-A3B</i> , gmu 0.75					
1	0.95	2.32	1.09	0.94	1.00
8	0.98	2.44	1.00	1.00	1.00
32	1.00	2.45	1.01	0.99	1.00
64	1.00	2.48	0.94	1.00	1.00
128	1.02	2.43	0.99	1.02	1.01
256	1.18	1.62	4.61	0.80	1.46
<i>Qwen3.5-35B-A3B</i> , gmu 0.9 (no knee)					
1	0.93	2.32	0.91	0.93	1.00
8	0.97	2.44	0.65	0.98	1.00
32	0.98	2.45	1.06	0.98	1.00
64	0.98	2.46	0.94	0.98	1.00
128	1.00	2.48	0.99	0.99	1.00
256	1.03	2.26	1.44	0.95	1.07
<i>Qwen3.5-397B-A17B</i> , gmu 0.9 (no knee)					
1	0.93	2.00	0.95	0.92	1.00
8	0.97	2.19	0.94	0.96	1.00
32	0.99	2.19	0.94	0.99	1.00
64	0.99	2.19	1.00	0.99	1.00
128	0.99	2.22	0.98	0.99	1.00
256	1.00	2.15	1.02	0.98	1.03

B Wide draft trees

Wider trees. Snapshotting cannot afford a wide tree; TreeWY can, and Table 2 in the main text is the primary width sweep. Table 4 below adds more shapes on the block-cost side of that trade: a store-all baseline’s per-request cost grows $10\times$ ($4 \rightarrow 40$ blocks) from $(1, 1, 1)$ to $(3, 3, 3)$, while TreeWY’s stays at one block regardless of w . Acceptance length keeps rising over that same range ($1.883 \rightarrow 2.786$) even though it is bounded by draft depth rather than width (nearly flat from $(2, 2, 2)$ to $(3, 3, 3)$ while blocks go $15 \rightarrow 40$): the gain comes from matching each level’s probability mass more often, not from a longer accepted path.

Table 4: More tree-width shapes (Qwen3.5-35B-A3B, depth-3, branching factor w ; primary width sweep in Table 2). “blk” is a store-all baseline’s per-request block cost ($N+1$); TreeWY holds one block at every w . Acceptance length matches the storeall reference within sampling noise.

shape	N	blk	storeall	treewy
$(1, 1, 1)$	3	4	1.979	1.883
$(2, 2, 2)$	14	15	2.854	2.776
$(3, 3, 3)$	39	40	2.868	2.786

C Admission vs. per-step cost

Reading the latency columns. Admitting more requests raises *per-request* latency at unchanged per-step cost, so a TPOT ratio below 1 is partly the price of that concurrency: across the 35B points the penalty tracks batch growth almost exactly (Pearson $r = -0.88$ against the running-batch ratio, $+0.88$ for throughput). At the 31 points where the baseline saturated, TreeWY pays $0.83\times$ on TPOT yet its mean *end-to-end* latency still improves ($1.17\times$), because requests the baseline could not admit stop waiting; at the other 74 its batch is matched to within 1% and TPOT is 0.98; that, not the aggregate, is our per-step cost (same decomposition for ReplaySSM, $r = -0.79$).

D ReplaySSM comparison

Comparison with ReplaySSM. ReplaySSM [8] (approach (ii)) is the closest concurrent point with a chain-mode implementation, so we ran the identical 35B sweep against it (Tables 5, 6 below). The two arms sit on different vLLM builds, so each is normalised to the snapshotting baseline from *its own* image; Table 5 lists both baselines so the drift is visible: the store-all throughputs agree to within 1%, so the two arms’ own-baseline ratios below are directly comparable to each other. Both delete the per-position snapshots and free the same order of KV headroom, TreeWY marginally more ($0.96\times$ ReplaySSM’s peak usage), and both preserve acceptance on real prompts (mean $|\Delta|$ acceptance length 0.029 vs. 0.037). ReplaySSM turns its headroom into more throughput on a *chain* ($1.12\text{--}1.20\times$ its baseline vs. our $0.99\text{--}1.08\times$) and preempts less often. That margin is not the solve (the solve is shared): our working hypothesis is that it is the scheduling around it, since they defer writing state on a verify step where we materialize the accepted state on every commit, but this is an attribution, not a validated claim: we have not prototyped a deferred-write commit path ourselves to check that it actually closes the gap. If it holds, it is an implementation lever open to us and orthogonal to the closed form; testing it, alongside the analogous sweep we have not yet run against Bole [12] (approach (iv), Section 2, which targets a different serving engine), is future work.

Table 5: Why the ratios differ, decomposed at 256 concurrency (Qwen3.5-35B-A3B, all 15 points: 5 workloads (ShareGPT, spec-bench, synthetic balanced-chat, generation-heavy, summarize-heavy) $\times$ 3 gmu; BurstGPT excluded, its own independent variable). *Absolute* levels first, then the identity. Both methods admit essentially the same work: 238 vs 224 peak running requests, and at gmu 0.6 per workload they are indistinguishable (253/255, 254/255, 256/256, 256/256, 109/107 for ShareGPT/spec-bench/balanced-chat/generation-heavy/ summarize-heavy), because both simply reach the offered concurrency once the snapshots are gone. TreeWY’s larger *batch ratio* in Table 6 is therefore a denominator artifact: its own baseline happened to admit slightly fewer (92 vs 97 on balanced-chat, 61 vs 63 on summarize-heavy). What separates the two is per-token cost, and once both are saturated and decode-bound the throughput ratio is just $\text{batch} \div \text{cost}$, which closes to within 3% for both. So at this operating point the memory mechanism is a tie and the whole difference is per-step efficiency.

	Ours		ReplaySSM	
	Store-all	TreeWY	Store-all	ReplaySSM
peak running requests	151.1	238.2	153.6	223.9
output tok/s	8469	10094	8250	11357
mean TPOT (ms)	16.4	21.2	17.0	17.6
p99 TTFT (ms)	8167	2746	8965	2958
<i>ratio vs own baseline, and the decode-bound identity</i>				
admitted batch	1.577		1.458	
per-token cost	$1.295\times$		$1.035\times$	
batch $\div$ cost (predicted)	1.218		1.408	
throughput (measured)	1.192		1.376	

Table 6: The same load axis for both methods (Qwen3.5-35B-A3B), each normalised to the store-all baseline from *its own* vLLM build, so the two columns are comparable even though the builds are not. *Per-token* speedup is based on mean TPOT. ReplaySSM leads on throughput at every load, and the gap is widest exactly where neither method is admitting anything extra (batch 1.00, loads 1–64), which is what identifies it as per-step cost rather than a memory effect. Past the knee the two converge on the memory-driven columns: at 256 concurrency the peak-KV reduction (1.56 vs. 1.57) and the admitted batch (1.58 vs. 1.46) are within a few percent, i.e. both free the same headroom and spend it the same way. On the BurstGPT arrival-rate sweep at 64 req/s the p99 TTFT improvements are $10.72\times$ and $10.20\times$ respectively. The per-token column is the clearest read on per-step cost: ReplaySSM leads it at every load, and Table 5 decomposes the 256-concurrency point.

concurrency	Output tput. gain $\uparrow$ ours / Replay	Peak KV reduction $\uparrow$ ours / Replay	p99 TTFT speedup $\uparrow$ ours / Replay	Per-token speedup $\uparrow$ ours / Replay
1	0.94 / 1.06	2.36 / 2.15	1.00 / 0.98	0.94 / 1.07
8	0.99 / 1.08	2.43 / 2.26	0.87 / 1.20	1.00 / 1.07
32	1.00 / 1.11	2.45 / 2.29	1.03 / 0.99	0.99 / 1.11
64	0.99 / 1.17	2.46 / 2.29	0.96 / 1.01	0.99 / 1.17
128	1.07 / 1.27	2.15 / 2.02	1.76 / 1.60	0.94 / 1.14
256	1.19 / 1.38	1.56 / 1.57	2.97 / 3.03	0.77 / 0.97

E Additional serving curves

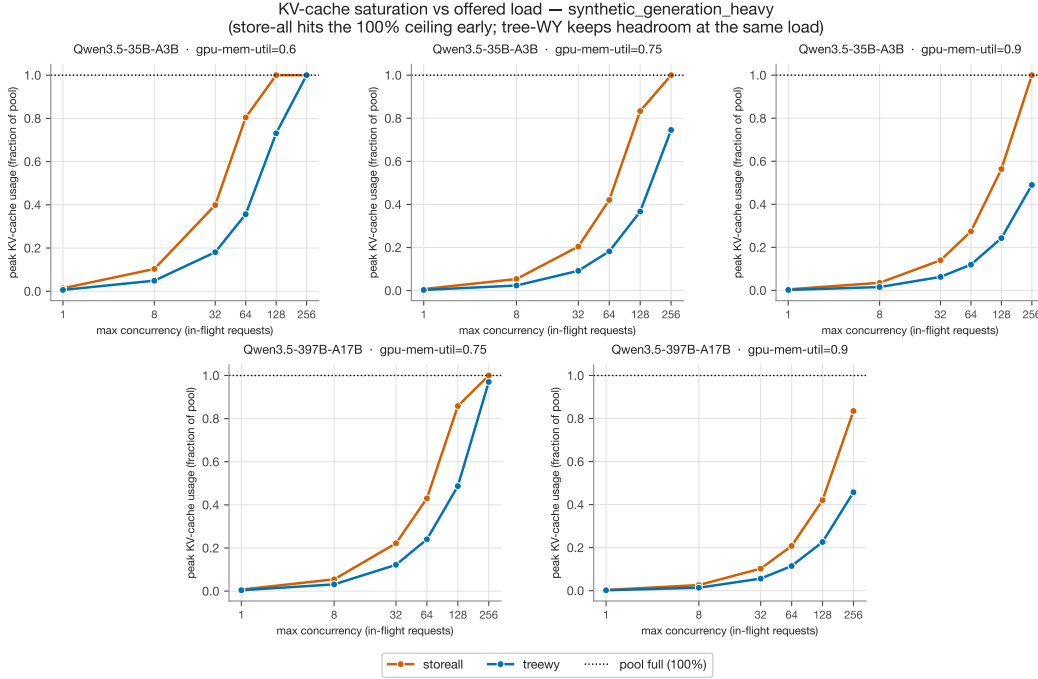


Figure 1: Peak KV-cache usage (fraction of pool) vs. offered load, per model and GPU memory utilization (gmu). The snapshotting baseline (storeall) hits the 100% ceiling early and must then queue the requests it cannot fit; TreeWY holds 2–3 $\times$ more headroom at the same load. That freed capacity buys the admission, throughput and TTFT wins (Fig. 2).

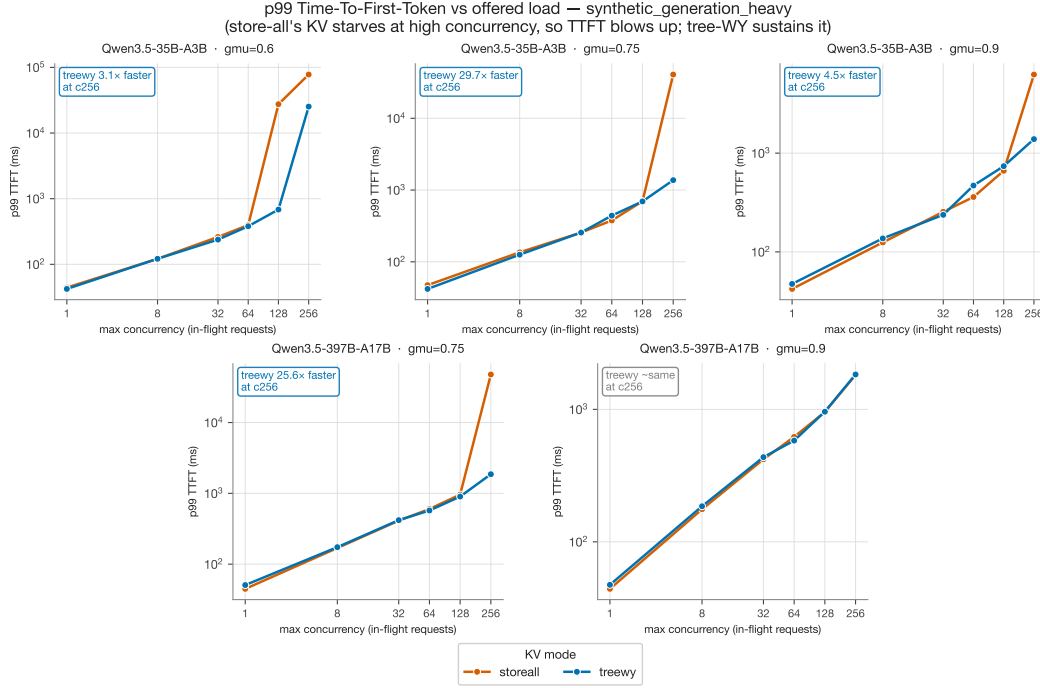


Figure 2: p99 time-to-first-token (TTFT) vs. offered load (generation-heavy workload). As concurrency rises, the snapshotting baseline’s KV pool saturates and it must queue requests, so its TTFT blows up (log scale): at 256 concurrency it is up to $\sim 30\times$ (35B, utilization 0.75) and $26\times$ (397B, utilization 0.75) worse than TreeWY. The largest gap over the whole sweep, $40\times$ (683 vs. 27489 ms), is at 128 concurrency (35B, utilization 0.6), where the baseline saturates but TreeWY still does not. The gap closes only when memory is slack (rightmost panels).

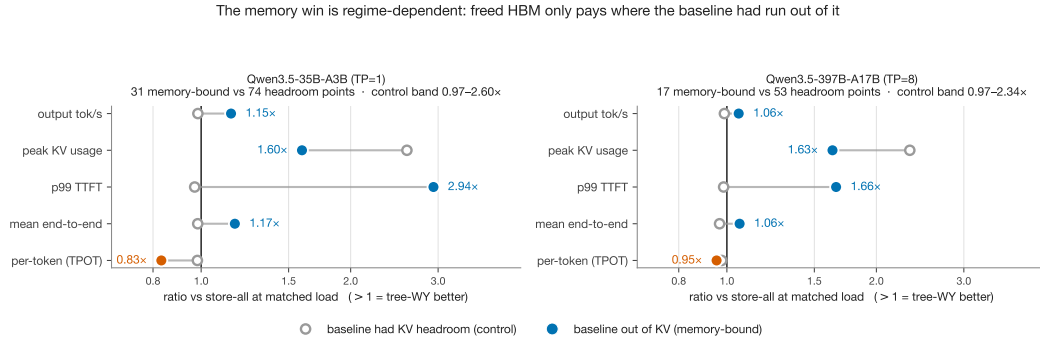


Figure 3: Every ratio in Table 1 split by whether the *baseline* had run out of KV headroom at that load, the partition behind Appendix C. Open markers are the headroom control: by construction every metric sits within a few percent of parity there. Filled markers are where the freed HBM is actually spent: at 35B (31/105 points memory-bound) TreeWY reaches $1.15\times$ throughput, $2.94\times$ lower p99 TTFT and $1.17\times$ lower mean end-to-end latency at $1.60\times$ more KV headroom, for a $0.83\times$ per-token cost; 397B (17/70 points memory-bound) shows the same shape at smaller magnitude ($1.06\times$ throughput, $1.66\times$ lower p99 TTFT, $1.63\times$ more KV headroom, $0.95\times$ per-token cost), since its sweep reaches fewer points past the knee. This is the single figure that makes the paper’s central qualifier legible at a glance: the win is not a fixed multiplier, it is conditional on memory pressure, and averaging over both regimes (as a single headline number would) hides exactly that.

F Compute resources

GPU-hours. All runs use B200 GPUs (178 GiB HBM/device). Each job serves one (model, mode, gmu) combination through its full 35-point concurrency/arrival-rate sweep in one continuous server session. The 35B main sweep (Table 1, Appendix A) is 6 jobs (`storeall/treewy` $\times$ 3 gmu) at TP1, ≈ 11.8 GPU-hours; the 397B main sweep is 4 jobs (`storeall/treewy` $\times$ 2 gmu) at TP8, ≈ 62.5 GPU-hours; the ReplaySSM comparison (Appendix D) is 6 further 35B/TP1 jobs, ≈ 11.1 GPU-hours. Total: ≈ 85 GPU-hours across the three reported sweeps. This is a lower bound: it excludes one-time model-loading/server-startup overhead at the start of each sweep (~ 1 h for the 397B model, ~ 30 min for 35B), which does not recur across jobs within the same sweep.