

A DEEPSEEK-R1 RESULTS

Table 6 presents end-to-end results for DeepSeek-R1 on MultihopRAG and NarrativeQA datasets, evaluated on $16 \times H20$ and $32 \times H20$ GPUs. ContextPilot achieves $1.81 \times$ and $1.52 \times$ throughput improvements on MultihopRAG and NarrativeQA respectively, with these speedups remaining consistent across both 16 and 32 GPU deployments using context-aware routing.

Table 6. DeepSeek-R1 results formatted for vertical readability. By stacking hardware configurations, the table width is minimized while preserving all data points.

Method	Hardware	Prefill TP (tok/s)	Cache Hit	F1 (%)
<i>Dataset: MultihopRAG</i>				
Vanilla	$16 \times H20$	9636.69	5.12%	64.15
	$32 \times H20$	18406.08	4.17%	64.15
ContextPilot w/o Annotations	$16 \times H20$	17498.75	60.37%	64.09
	$32 \times H20$	33072.64	58.41%	64.09
ContextPilot (Ours)	$16 \times H20$	17498.75	60.37%	64.68
	$32 \times H20$	33072.64	58.41%	64.68
<i>Dataset: NarrativeQA</i>				
Vanilla	$16 \times H20$	8687.98	6.08%	40.20
	$32 \times H20$	16247.32	5.14%	40.20
ContextPilot w/o Annotations	$16 \times H20$	13201.52	38.24%	40.38
	$32 \times H20$	24686.84	35.71%	40.38
ContextPilot (Ours)	$16 \times H20$	13201.52	38.24%	41.08
	$32 \times H20$	24686.84	35.71%	41.08

B ATTENTION MAP ANALYSIS

Since context engineering (Rajasekaran et al., 2025) and in-context learning (Kirsch et al., 2022) strongly influence model inference, we analyze attention patterns when explicit annotations are introduced to recall the original relevance ranking.

Figure 9 and Figure 10 compare the final-layer attention maps of Qwen3 and LLaMA3.3 under this setup. When given explicit document-priority cues, both models exhibit consistent attention behaviors despite architectural differences. They correctly focus on document tokens ([Doc.1], [Doc.2], [Doc.3]), reflecting awareness of the mismatch between the aligned and original sequences, as indicated by intersections between queries in the annotation region and keys in the context region. As the context re-aligns with the original sequence, both models emphasize ([Doc.2]) while parsing ([Doc.1]) and ([Doc.3]), showing that the cue ([Doc.2] > [Doc.1] > [Doc.3]) effectively directs cross-document attention. Hence, explicit annotations reshape internal attention, aligning it with semantic rather than positional priority.

This finding supports a central hypothesis of ContextPilot: explicit annotations can make a comeback from the accuracy lost to alignment and filtering by re-establishing alignment with the original retrieval semantics.

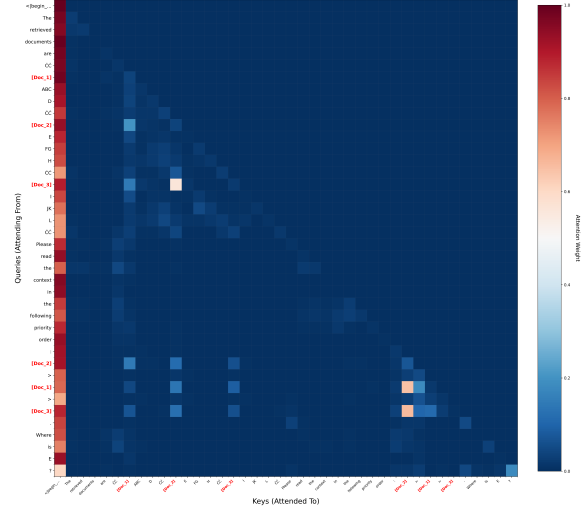


Figure 9. Attention map of the last layer attention of LLaMA3.3 (Head 14) for the prompt “The retrieved documents are [Doc.1] ABCD [Doc.2] EFGH [Doc.3] IJKL. Please read the context in the following priority order: [Doc.2] > [Doc.1] > [Doc.3]. Where is E?”.

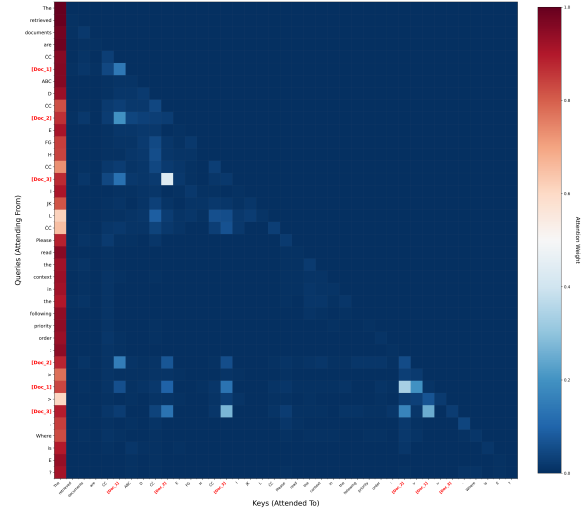


Figure 10. Attention map of the last layer attention of Qwen3 (Head 9) for the prompt “The retrieved documents are [Doc.1] ABCD [Doc.2] EFGH [Doc.3] IJKL. Please read the context in the following priority order: [Doc.2] > [Doc.1] > [Doc.3]. Where is E?”.

C DOCUMENT ACCESS DISTRIBUTION

Figure 11 shows the cumulative distribution of document access frequency across three RAG datasets. A small fraction of documents accounts for the majority of retrievals: the top

20% most frequently accessed documents cover 79.2% of retrieval events on MultihopRAG, 57.4% on NarrativeQA, and 49.6% on QASPER. This heavy-tailed distribution confirms that real-world RAG workloads exhibit substantial context overlap across sessions, motivating context reuse through prefix-aligned alignment.

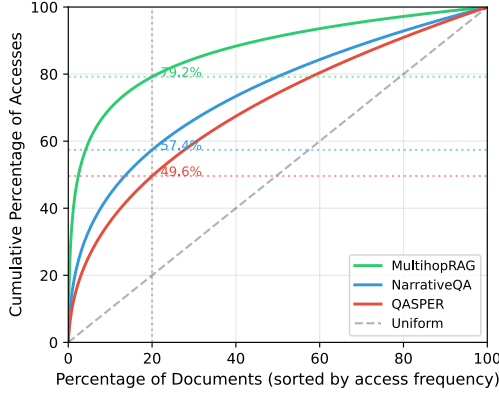


Figure 11. Document access distribution (CDF) across three datasets. The vertical dashed line marks 20% of documents; horizontal lines indicate the cumulative percentage of accesses from these top documents.

D ADDITIONAL EVALUATION DETAILS

D.1 Time-Series Metrics

Figure 12 shows how cache hit ratio evolves as the workload progresses. ContextPilot maintains approximately 34% cache hit ratio compared to baseline’s 7% throughout the entire workload, demonstrating a sustained $5\times$ improvement that is not a transient warm-up effect.

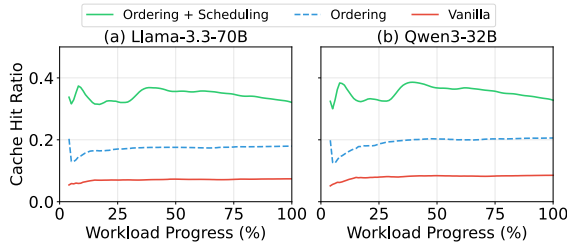


Figure 12. Cache hit ratio over workload progress for Llama-3.3-70B and Qwen3-32B. ContextPilot maintains $\sim 5\times$ higher hit ratio throughout execution.

Figure 13 presents cumulative cached tokens as a metric for radix tree prefix reuse. ContextPilot achieves 10.33M cached tokens versus baseline’s 2.42M at completion on Llama3.3-70B-Instruct ($4.27\times$), and 10.50M versus 2.75M

on Qwen3-32B ($3.82\times$). The “w/o Scheduling” variant ($6.85M$, $2.83\times$) confirms that both alignment and scheduling contribute to the improvement.

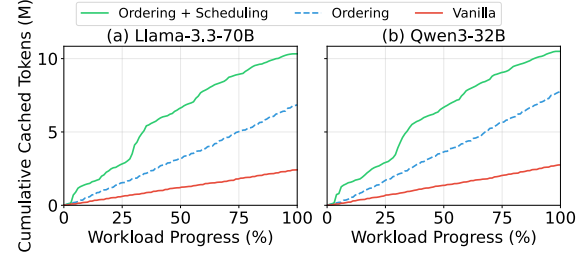


Figure 13. Cumulative cached tokens (radix tree reuse) over workload progress for Llama-3.3-70B and Qwen3-32B. ContextPilot achieves $4\times$ better prefix reuse.

D.2 Accuracy Breakdown

Table 7 provides a detailed breakdown of accuracy contributions from each ContextPilot component.

Table 7. Accuracy breakdown by component. Alignment alone: $\leq 1\%$ F1 variation; adding annotations: $+1.4$ – 4.4% gains over aligned baseline.

Model	Configuration	MultihopRAG	NarrativeQA
Qwen3-32B	Baseline	60.4%	28.4%
	+ Alignment	60.0%	28.2%
	+ Annotation	64.4%	29.6%
	+ Scheduling	64.4%	29.6%
Qwen3-4B	Baseline	35.2%	16.0%
	+ Alignment	34.5%	15.2%
	+ Annotation	36.6%	17.3%
	+ Scheduling	36.6%	17.3%

D.3 Per-Request Overhead

Table 8 reports the per-request overhead of ContextPilot components, measured on 2K requests with $k = 15$ on NVIDIA A6000.

Table 8. Per-request overhead breakdown. Total overhead (~ 0.7 ms) is negligible compared to seconds of prefill latency.

Component	Latency (ms)
Search	0.068
Alignment	0.047
De-duplication	0.600
Total	~ 0.7

E OPENCLAW PIPELINE DIAGRAM

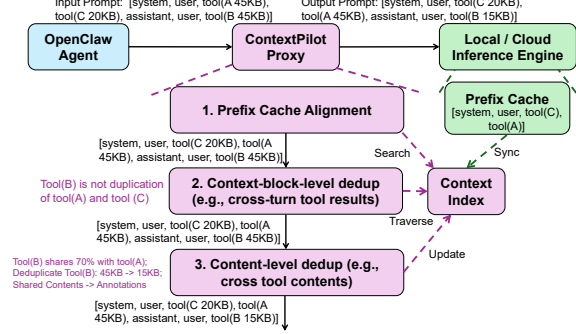


Figure 14. The OpenClaw + ContextPilot pipeline. The OpenClaw agent sends requests to the ContextPilot proxy, which aligns prompts with the prefix cache of the inference engine, deduplicates overlapping context both across turns and within context blocks, and forwards optimized prompts to the inference engine.

F SYSTEM OVERHEAD WITH ZERO CONTEXT OVERLAP

Zero context overlap represents the worst case for ContextPilot, as it isolates pure system overhead with no reuse benefit. Using a synthetic RAG workload with no retrieval overlap, ContextPilot adds only 0.72 s of prefill latency for 1K contexts (one-hour job), demonstrating that querying the context index during operation incurs negligible overhead.

G IMPACT OF PREFIX CACHE SIZE

The prefix KV-cache stores precomputed attention states for reuse across requests; a larger cache retains more contexts, increasing the chance of a cache hit and reducing redundant prefill computation. We evaluate this effect on MultihopRAG across A6000 (48 GB) and H100 (80 GB) GPUs. Because ContextPilot aligns contexts with the prefix cache to maximize overlap, it benefits disproportionately from larger caches: SGLang hit ratios improve from 29.64% to 33.97%, and vLLM from 35.90% to 43.4%, translating directly into higher prefill throughput. In contrast, baselines see smaller gains since their context alignment does not systematically exploit the additional capacity.

H DETAILED ALGORITHM PSEUDOCODE

This appendix collects all algorithm pseudocode referenced in the main text. Algorithm 1 details the context index tree search, Algorithm 2 describes context alignment for prefix sharing, Algorithm 3 formalizes context de-duplication, Algorithm 4 provides the full pseudocode for context index construction via hierarchical clustering, and Algorithm 5

details the tree-based request grouping and scheduling procedure.

Algorithm 1 Context Index Tree Search

Require: Context $C = \{b_1, \dots, b_k\}$, root node R

Ensure: Search path, best matching node

```

1:  $cur \leftarrow R, path \leftarrow []$ 
2: while  $cur$  has children do
3:    $best \leftarrow \arg \min_{c \in cur.children, C \cap c.blocks \neq \emptyset} DIST(C, c)$ 
4:   if no overlapping child found then
5:     break
6:   end if
7:    $path.APPEND(\text{index of } best)$ 
8:   if  $best$  is leaf then
9:     return  $path, best$ 
10:  end if
11:   $cur \leftarrow best$ 
12: end while
13: return  $path, cur$ 

```

Algorithm 2 Context Alignment for Prefix Sharing

Require: Context C with context blocks, Context Tree root node R

Ensure: Aligned context C'

```

1:  $bestMatch \leftarrow \text{FINDBESTMATCHNODE}(C, R)$ 
2: if  $bestMatch = \text{null}$  then
3:    $C' \leftarrow C$ 
4:   return  $C'$ 
5: end if
6:  $prefix \leftarrow bestMatch.context$ 
7:  $remaining \leftarrow C \setminus prefix$ 
8:  $C' \leftarrow prefix \cup remaining$ 
9: return  $C'$ 

10: function  $\text{FINDBESTMATCHNODE}(C, R)$ 
11: if  $C$  is initialization context then
12:   return  $C.parent$ 
13: else
14:   return  $\text{SEARCHTREE}(C, R)$ 
15: end if
16: end function

```

Algorithm 3 Context De-duplication

Require: New context C_{new} , context index $\mathcal{I}$, conversation ID id , chunk modulus M

Ensure: Deduplicated context C' , location annotations H

- 1: $S \leftarrow \mathcal{I}.\text{seen_blocks}[id]$ $\triangleright$ blocks indexed in prior turns
- 2: $V \leftarrow \mathcal{I}.\text{seen_subblocks}[id]$ $\triangleright$ sub-block hashes from prior turns
- 3: $C' \leftarrow [], H \leftarrow []$
- 4: **for** each block b in C_{new} **do**
- 5: **if** $b \in S$ **then**
- 6: $h \leftarrow$ “refer to $[b]$ in previous conversation”
- 7: $C'.\text{APPEND}(h); H.\text{APPEND}(h)$
- 8: **else**
- 9: $\{s_1, \dots, s_m\} \leftarrow \text{CDC}(b, M)$ $\triangleright$ content-defined chunking
- 10: **for** $j = 1$ to m **do**
- 11: $f \leftarrow \text{HASH}(s_j)$
- 12: **if** $f \in V$ **and** $V[f] \neq b$ **then**
- 13: Replace s_j with location annotation to $V[f]$
- 14: **else**
- 15: $V[f] \leftarrow b$
- 16: **end if**
- 17: **end for**
- 18: $C'.\text{APPEND}(\text{CONCAT}(s_1, \dots, s_m))$
- 19: **end if**
- 20: **end for**
- 21: $S \leftarrow S \cup C_{\text{new}}$ $\triangleright$ register for future turns
- 22: $\mathcal{I}.\text{seen_subblocks}[id] \leftarrow V$ $\triangleright$ persist sub-block hashes
- 23: **return** C', H

Algorithm 4 Context Index Construction via Hierarchical Clustering

Input: Batch of n contexts $\mathcal{S} = \{s_1, s_2, \dots, s_n\}$, distance function $d(\cdot, \cdot)$

- 1: **// Phase 1: Distance computation and clustering**
- 2: $D \leftarrow$ pairwise distances using $d(s_i, s_j) \triangleright$ GPU or CPU
- 3: $Z \leftarrow \text{LINKAGE}(D)$ $\triangleright$ hierarchical clustering
- 4: **// Phase 2: Build tree with deduplication**
- 5: **for** $i = 1$ to n **do**
- 6: Create leaf node v_i with $v_i.\text{content} \leftarrow s_i$
- 7: **if** $\exists v_j$ with $v_j.\text{content} = s_i$ **then**
- 8: Redirect $v_i \rightarrow v_j; v_j.\text{freq} += 1$ $\triangleright$ dedup
- 9: **end if**
- 10: **end for**
- 11: **for** each merge (c_1, c_2, δ) in Z **do**
- 12: $v \leftarrow$ new internal node
- 13: $v.\text{content} \leftarrow c_1.\text{content} \cap c_2.\text{content}$
- 14: $v.\text{children} \leftarrow [c_1, c_2]; c_1.\text{parent}, c_2.\text{parent} \leftarrow v$
- 15: **end for**
- 16: Remove empty internal nodes; relink children to grand-parents
- 17: **// Phase 3: Top-down prefix alignment**
- 18: Compute search paths from root to all nodes
- 19: **for** each node v in BFS order from root **do**
- 20: **if** v is not root **and** $v.\text{parent}.\text{docs} \neq \emptyset$ **then**
- 21: $v.\text{docs} \leftarrow v.\text{parent}.\text{docs} \oplus (v.\text{docs} \setminus v.\text{parent}.\text{docs})$
- 22: **end if**
- 23: **end for**
- 24: **return** aligned contexts from leaf nodes, search paths

Algorithm 5 Search-Path-Based Request Grouping and Scheduling

- 1: **Input:** N contexts with search paths $\{P_1, \dots, P_N\}$ from alignment
- 2: **Output:** Scheduled execution order
- 3: **// Phase 1: Group by root prefix** $\triangleright O(N)$
- 4: $\mathcal{G} \leftarrow$ empty map
- 5: **for** $i = 1$ to N **do**
- 6: $key \leftarrow P_i[0]$ $\triangleright$ first element of search path
- 7: $\mathcal{G}[key].\text{APPEND}(i)$
- 8: **end for**
- 9: **// Phase 2: Sort within each group** $\triangleright O(N \log N)$
- 10: **for** each group $G \in \mathcal{G}$ **do**
- 11: Sort G by $|P_i|$ descending $\triangleright$ longer paths first
- 12: **end for**
- 13: **// Phase 3: Order groups and flatten**
- 14: Sort groups by $|G|$ descending $\triangleright$ largest groups first
- 15: **return** concatenation of all groups
