



KVSWAP: Disk-aware KV Cache Offloading for Long-Context On-device Inference

Huawei Zhang
schz@leeds.ac.uk
University of Leeds

Chunwei Xia
C.Xia@leeds.ac.uk
University of Leeds

Zheng Wang
Z.Wang5@leeds.ac.uk
University of Leeds

Abstract

Language models (LMs) are enabling a new class of mobile and embedded AI applications, including meeting and video summarization, document analysis, and other multi-input workloads that require long-context reasoning. Running these models locally improves privacy, enables offline use, and lowers serving cost. However, on-device long-context inference is fundamentally constrained by a *memory capacity wall*: the key-value (KV) cache grows linearly with context length and batch size, quickly exceeding available memory. Existing KV-cache offloading schemes are designed to transfer cache data from GPU memory to CPU memory; however, they are not suitable for embedded and mobile systems, where the CPU and GPU (or NPU) typically share unified memory, and secondary storage offers limited, highly asymmetric I/O bandwidth. We present KVSWAP, the first KV-cache management framework for resource-constrained mobile systems that uses storage, rather than CPU RAM, as the KV-cache backing store for long-context inference. KVSWAP stores the full cache on disk, uses compact in-memory metadata to predict and preload future accesses, overlaps computation with hardware-aware storage transfers, and optimizes read patterns for device characteristics. Our evaluation shows that across representative LMs and storage types, KVSWAP delivers higher throughput under tight memory budgets while maintaining generation quality over existing KV cache offloading schemes.

CCS Concepts

• **Human-centered computing** → **Ubiquitous and mobile computing**.

Keywords

Large Language Models, Long Context Inference, Mobile System

ACM Reference Format:

Huawei Zhang, Chunwei Xia, and Zheng Wang. 2026. KVSWAP: Disk-aware KV Cache Offloading for Long-Context On-device Inference. In *The 24th Annual International Conference on Mobile Systems, Applications and Services (MobiSys '26)*, June 21–25, 2026, Cambridge, United Kingdom. ACM, New York, NY, USA, 16 pages. <https://doi.org/10.1145/3745756.3809234>

1 Introduction

Language models (LMs) have shown strong capabilities in long-form document analysis and multimedia understanding [40, 78, 83]. These underpin emerging mobile and embedded applications such

as smart note-taking for meetings [24, 53], document analysis [9, 27], voice assistants for transcription and summarization [13, 29], and video search with natural language queries [13, 28]. Increasingly, users expect these capabilities to run directly on their devices.

On-device deployment offers clear benefits: enhanced privacy by keeping data on-device, offline functionality under poor connectivity, and avoiding LM API costs. With the growing availability of powerful GPUs and NPUs on modern system-on-chips [3, 7, 56], local deployment is becoming feasible.

However, running LMs efficiently on resource-constrained devices like mobile and embedded AI systems remains challenging due to limited memory capacity and I/O bandwidth. While quantization [26, 44] and parameter offloading [11, 73] reduce model footprint, a critical bottleneck lies in managing the key-value (KV) cache. Unlike fixed model weights, the KV cache grows linearly with sequence length and batch size, often surpassing model size (as shown in Fig. 1). On a 4B-parameter LM, batched inputs over 16K tokens already push the KV cache beyond device memory, causing latency spikes or out-of-memory failures. As LMs are increasingly applied to long-context tasks, practical on-device inference requires efficient KV cache management.

Prior work reduces KV cache overhead by offloading from GPU to CPU memory [41, 60], mitigating GPU pressure in server settings with abundant CPU RAM and high memory bandwidths. These techniques, however, are ill-suited for mobile and embedded systems, where GPUs or NPUs typically share only 8-32 GB of RAM with the CPU. A natural alternative is to offload the KV cache to non-volatile storage¹, such as NVMe-, UFS-, or eMMC-based devices. However, without careful optimization, this approach severely degrades performance: mobile memory bandwidth is roughly 100 GB/s, while disk bandwidth is often just 1 GB/s - two to three orders of magnitude lower than their server counterparts.

Can we offload the KV cache to disk for long-context, on-device inference without sacrificing throughput or generation quality? In answer, we propose KVSWAP, a new framework for extended-context LM inference on resource-constrained devices. Built on FlexGen [58], KVSWAP supports diverse LM architectures and storage types, with simple APIs to use. KVSWAP offloads KV cache entries to disk and partially reloads them during generation. It stores the complete KV cache on disk but maintains a highly compact in-memory K cache representation to predict which entries on the disk are critical for each layer. These entries are prefetched into a buffer ahead of computation, reducing disk stalls without compromising generation quality. The KVSWAP runtime automatically overlaps disk transfers with compute and synchronizes memory-disk copies.



This work is licensed under a Creative Commons Attribution 4.0 International License. *MobiSys '26, Cambridge, United Kingdom*
© 2026 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-2027-7/2026/06
<https://doi.org/10.1145/3745756.3809234>

¹We use *disk* for non-volatile storage and *memory* for main memory.

The design of KVSWARE needs to overcome two major challenges introduced by on-device disk offloading. First, the ever-increasing KV-cache size introduced by long-context is a challenge for resource-constrained mobile and embedded devices. The limited memory capacity and strict per-application memory quotas result in a *tight memory budget*, far beyond what existing KV-cache offloading methods, which were primarily designed for server-side deployments, can accommodate. To address this, KVSWARE proposes a compression and reconstruction algorithm that enables aggressive memory compression while still achieving information recovery with adequate quality.

Second, embedded storage devices based on NVMe, UFS, and eMMC often exhibit *read amplification*, where the controller reads larger physical units (e.g., entire NAND pages) than requested [30, 52], resulting in extra data movement and unnecessary read operations. To mitigate this, KVSWARE groups KV entries at appropriate granularities and optimizes access patterns. It integrates a software cache (reuse buffer) to retain recently accessed entries across generations. This reduces redundant I/O, which is particularly important for low-bandwidth devices (< 200 MB/s) and batched inference, where data movement dominates latency. Together, these specified designs and optimizations enable KVSWARE for efficient long-context on-device inference with negligible loss of generation quality.

Recent work explores storing the full KV cache on disk with selective retrieval [19, 46, 76]. Designed for cloud serving, these methods mainly optimize time-to-first-token (TTFT) to reduce start-up latency. However, they are ineffective for iterative decoding on resource-constrained devices, as they still incur high memory overhead and assume static rather than evolving KV caches. KVSWARE addresses this gap by targeting *the decoding stage of LM inference*, efficiently managing dynamic KV caches to reduce memory use while preserving generation quality. KVSWARE also differs from GPU-CPU offloading and sparse-attention systems designed for server systems [17, 41, 58–61, 84]. It is the first framework built for KV-cache offloading under the tight memory and I/O constraints of mobile devices, enabling high-quality long-context on-device generation for single-request and batched workloads like meeting summarization and batch video processing.

We evaluate KVSWARE on text, reasoning and video LMs across model sizes, memory budgets, disk types, batch sizes, and context lengths. Compared with the strongest KV cache offloading baseline, KVSWARE delivers up to 1.8× (NVMe) and 4.1× (eMMC) throughput improvements at 32K context length under the same memory budget, while providing *substantially higher* generation quality. Against the industry-grade vLLM [39] working in an idealized setting where all available memory is dedicated to the KV cache - KVSWARE achieves comparable or higher throughput but using 11.0× less KV cache memory.

This paper makes the following contributions. It presents:

- the first disk-based KV cache offloading framework for efficient long-context, on-device LM decoding;
- a memory-efficient algorithm and system co-design that accurately identifies and accesses groups of critical KV entries, achieving substantial memory compression while preserving quality;
- a runtime system for managing and caching KV entries, with optimized computation pipeline and reduced I/O transfers.

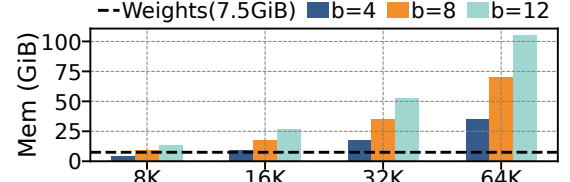


Figure 1: KV cache memory footprint across batch sizes (b) and sequence lengths (x-axis) of Qwen3-4B.

Online materials. The code and data associated with this work are available at <https://github.com/hwhwz23/KVSWARE-CODE.git>.

2 Background and Motivation

2.1 Language Models

Language models (LMs) process queries through a two-stage pipeline: *prefilling* and *decoding*. During prefilling, the model reads the entire input prompt to generate the first output token. During decoding, the model produces subsequent tokens one at a time, each conditioned on the previously generated tokens. Prefilling happens only once, whereas decoding is repeated for every output token. For open-ended tasks like document summarization or dialogue, decoding can involve *hundreds or even thousands* of steps, and this becomes even more common when using reasoning-focused models with chain-of-thought (CoT) [68].

Attention is central to LMs, where each token is represented by a *Query* (Q), *Key* (K), and *Value* (V) vector. Scores are computed by comparing queries with keys, and the resulting weights are applied to values to form contextualized representations. *Multi-head attention* (MHA) [65] extends standard attention by computing multiple sets of QKV vectors in parallel, but increases memory use since each head stores its own keys and values. *Grouped-query attention* (GQA)[10] addresses this by sharing keys and values across heads, reducing memory overhead while maintaining accuracy. GQA is widely adopted in state-of-the-art models, including Qwen3 [74].

During generation, LMs use a *key-value (KV) cache* to store keys and values from past steps, enabling the reuse of computations and reducing inference latency [55].

2.2 On-device KV Cache Memory Wall

The KV cache serves as a critical component of program state in on-device applications of LMs and is therefore usually required to be memory-resident. Fig. 1 shows the KV cache memory footprint for the Qwen3-4B model (W16A16) under varying batch sizes and context lengths. The model weights alone occupy 7.5 GiB of RAM. At a context length of 16 K and batch size 4, the KV cache reaches 9 GiB, already exceeding the weight size. Because KV cache size scales linearly with both batch size and context length, it grows rapidly: at 32 K context and batch size 12, it reaches 54 GiB.

On embedded and mobile systems with limited memory capacity (e.g., 8 GiB) and strict per-application quotas (to avoid disrupting other applications and OS services), retaining the full KV cache in memory becomes infeasible. This results in a *KV cache memory capacity wall*, where the KV cache overtakes model weights as the dominant memory consumer, limiting the scalability of long-context on-device inference.

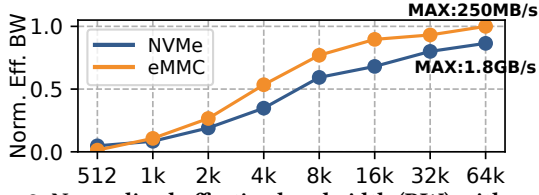


Figure 2: Normalized effective bandwidth (BW) with varying block sizes (bytes).

2.3 KV Cache Disk Offloading

Recent studies have shown that only a small fraction of tokens contribute significantly to attention score computation [42, 84]. This suggests that retaining a compact set of influential KV cache entries in memory can substantially reduce memory usage without compromising model accuracy. However, during iterative decoding, the set of influential tokens evolves with the context and query [41, 60, 61]. As a result, maintaining a static subset of KV entries becomes ineffective. This motivates our approach to dynamically load KV cache entries by predicting the most relevant tokens at each generation step.

Existing techniques improve KV cache efficiency by offloading data from GPU to CPU memory [41, 60], where both memory tiers offer high bandwidth and low latency. In contrast, our work targets embedded and mobile platforms, where RAM is limited and subject to strict constraints. We instead offload KV cache from memory to disk, operating under a much more constrained storage hierarchy with significantly lower bandwidth and higher latency.

To assess the feasibility of disk offloading, we profile the random read bandwidth of two representative types of disks - NVMe² and eMMC - under varying block sizes. We focus on random, rather than sequential access patterns because our setting requires *selectively* retrieving critical KV entries, operations that inherently produce predominantly random accesses. From Fig. 2, we observe:

Diverse disk I/O bandwidths. Disk types vary significantly in I/O bandwidth: NVMe can reach up to 1.8 GB/s, while eMMC can be limited to around 250 MB/s.

Poor I/O utilization for small KV entries. All storage types show severe bandwidth under-utilization for small transfer sizes. At 512 bytes, the typical size of a KV entry³, the effective bandwidth drops below 6% of peak bandwidth for both devices.

We observe that bandwidth utilization improves substantially with larger block sizes. KVSWAP addresses this by adopting a quality-aware, *group-wise* strategy that packs multiple KV entries into a single transfer. This amortizes I/O overhead and reduces fragmentation while incurring negligible model accuracy degradation (Sec. 3.3).

2.4 Drawbacks of Prior Offloading Schemes

High memory overhead. Recent KV cache offloading methods [41, 58, 60] are designed for CPU-GPU transfer on servers and are ill-suited for disk-based on-device offloading. They do not explicitly consider extremely tight memory constraints (e.g., <500 MiB), a regime that closely matches the practical application memory budget in common on-device deployment scenarios. Fig. 3a compares

²UFS-based storage is not supported on our platform, but it exhibits I/O bandwidth and characteristics similar to NVMe.

³128 (head dimension) × 2 (key and value) × 2 Bytes (float16) = 512 Bytes.

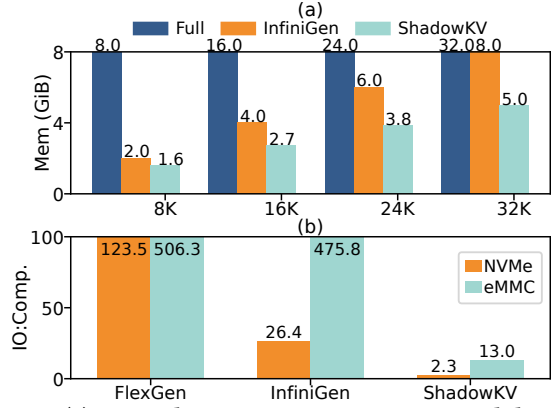


Figure 3: (a) KV cache management memory with batch=8 and varying context lengths (x-axis); (b) Decoding latency ratios of I/O to compute under 32K context length and batch=8.

the KV cache management memory of InfiniGen [41] and ShadowKV [60] against storing the full cache in memory for Llama3-8B. Their designs introduce a significant memory footprint, particularly for long contexts. For example, with a 16K context length and batch size 8, memory consumption already reaches 4 GiB and 2.7 GiB. Here, we attempted to reconfigure their parameters for tighter memory budgets and extend them to disk-based offloading, but this severely degrades generation quality (Sec. 5.1). We attribute this to the inability of their memory compression algorithms to tolerate high compression ratios. A detailed discussion is provided in Sec. 3.2 and Sec. 3.3.

Low I/O efficiency. The fine-grained head- or token-level operations create fragmented disk access patterns, causing excessive small disk reads and degrading I/O efficiency. As an example, Fig. 3b shows the decoding latency ratio of I/O to compute for FlexGen [58], InfiniGen [41], and ShadowKV [60]. Note that FlexGen does not select KV entries; instead, it loads the full KV cache into memory. All these methods exhibit ratios far exceeding 1, with some cases reaching over 100. While ShadowKV achieves the lowest ratios, it still reaches 13.0 on eMMC and 2.3 on NVMe. These high ratios indicate a severely imbalanced I/O–compute pipeline caused by inefficient I/O access. End-to-end results in Sec. 5.2 further validate the resulting low system throughput for these methods.

Lessons learnt. For on-device disk offloading, existing KV cache offloading methods suffer from *high memory overhead*, which limits scalability to longer contexts or larger batch sizes. In addition, their *low I/O efficiency* significantly limits system decoding, resulting in low generation speed and energy waste. KVSWAP is designed to avoid these pitfalls.

3 Our Approach

KVSWAP is a Python framework that supports PyTorch-based Transformer LMs and provides simple APIs for model serving and parameter tuning. Fig. 4 shows example usage for LM inference: an offline parameter tuning API (Fig. 4a) generates configurations used by the KVSWAP runtime (Fig. 4b). Using KVSWAP is straightforward and typically requires only a few dozen lines of Python code.

Fig. 5 gives a high-level overview of KVSWAP. KVSWAP reduces memory pressure and improves system efficiency by: (a) offloading


```
import KVSWARE
KVSWARE.Parameter_Tuning(
    model=model_obj, max_context_len=32768,
    max_batch_size=8, max_kv_mem=2200 # in MiB
).save("config.json")

(a) Offline parameter tuning.

engine = KVSWARE.Init( model=model_obj,
    ↪ kv_offload_path="...", config_file="config.json" )
outputs = engine.generate( inputs=[...],
    ↪ sampling_params=[...] )

(b) Model serving.
```

Figure 4: Simplified examples for using KVSWARE for offline parameter tuning (a) and model serving (b).

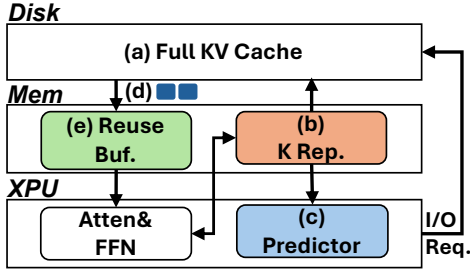


Figure 5: High-level overview of KVSWARE.

the full KV cache to disk; (b) maintaining a compact in-memory key cache representation to identify important KV entries using a predictor (c), which will be preloaded from the disk for decoding; (d) minimizing disk I/O overhead by structuring disk access patterns; and (e) effectively caching and reusing previously accessed KV entries through a reuse buffer. The KVSWARE runtime coordinates I/O and compute, tracks reusable entries across layers, fetches missing ones from disk for the attention kernel, and periodically updates the on-disk full cache and the in-memory representation with newly generated data. For shorter context, KVSWARE keeps the KV cache fully in memory when possible to minimize the data movement overhead.

Key technical contributions of KVSWARE include low-rank key cache summaries for importance prediction, and disk-aware grouped KV prefetching with buffer reuse. The predictor is detailed in Sec. 3.1 and Sec. 3.3, the in-memory compression scheme in Sec. 3.2, and the runtime scheduling and reuse buffer in Sec. 3.4.

3.1 Predicting Important KV Cache Entries

The KVSWARE KV predictor identifies, *at runtime*, a subset of KV entries needed for each attention computation. Like InfiniGen [41], it estimates the next layer’s attention scores from the current layer’s input, treating it as an approximate query (Q). These scores quantify the importance of each key and its associated value. Unlike [41], our setting assumes the entire KV cache resides on disk, where memory and I/O bandwidth are limited. To make prediction feasible, KVSWARE maintains a *compact* in-memory representation of the K cache. This compressed cache is used to compute approximate attention scores and select a configurable number of top-ranked KV entries *as groups* for preloading.

Thus, the predictor balances two objectives: (1) minimizing memory footprint, through K cache compression and reconstruction, and (2) improving disk efficiency, through grouped critical KV prediction. These two techniques are detailed in the following subsections.

3.2 Compressing K Cache

To enable prediction without an in-memory full K cache, KVSWARE maintains a *compressed* K cache in memory using low-rank approximation. Instead of storing every detail of the K matrix, we keep a smaller “summary” that captures its most important patterns. The compressed K cache is automatically managed and updated by the KVSWARE runtime (Sec. 3.4.1). Since the low-rank K cache is used only for estimating attention scores - not for actual attention computation - precision can be flexibly traded for memory efficiency.

We consider two designs: (a) *per-head compression*, compressing each head with a separate projection matrix, and (b) *joint-head compression*, merging head and embedding dimensions with a single projection. We adopt the latter for its unified representation and lower memory cost, reconstructing the multi-head structure during prediction (Sec. 3.3).

We first reshape the K cache into a matrix of size $N \times (H_k \cdot d)$, where N is the token count, H_k is the number of heads, and d is the head dimension. Each vector is then projected into a lower dimension r using a *pre-computed low-rank adapter* $A \in \mathbb{R}^{(H_k \cdot d) \times r}$, with $r \ll H_k \cdot d$. To obtain A , KVSWARE applies *Singular Value Decomposition* (SVD) [23] to a flattened K cache ($K_{\text{fltn}} \in \mathbb{R}^{N \times (H_k \cdot d)}$) collected during offline parameter tuning. By default, samples of this K cache are drawn from general-purpose datasets [51, 57], though users may provide their own through the KVSWARE API. Formally, $\text{SVD}(K_{\text{fltn}}) = U \text{diag}(S) V^T$, where $V \in \mathbb{R}^{(H_k \cdot d) \times m}$ and $m = \min(N, H_k \cdot d)$. The top- r right singular vectors of V form A , and the compressed K cache is $K_{\text{lr}} = \text{Flatten}(K)A$, with compression level controlled by the compression ratio, $\sigma = \frac{H_k \cdot d}{r}$. We found that the resulting low-rank adapter generalizes well to different datasets.

Compared with prior K cache compression scheme. While ShadowKV [60] also employs a low-rank K cache, its design and purpose differ from ours. ShadowKV compresses the K cache to reduce its memory footprint, but must reconstruct the full K cache for computation. Hence, it has to use a conservative compression ratio to limit accuracy loss. In contrast, KVSWARE uses the low-rank K cache only to predict critical KV entry indices, making it far more resistant to aggressive compression. Combined with our prediction algorithm (Sec. 3.3), this enables KVSWARE to maintain generation quality even under high compression ratios, whereas ShadowKV degrades significantly (Sec. 5.1). The construction method also differs. ShadowKV performs an online SVD at runtime, adding substantial prefill latency (e.g., 4.9× slower at 8K context length).⁴ KVSWARE instead precomputes the low-rank adapter offline, incurring no extra prefill cost.

3.3 Grouped Critical KV Entries Prediction

To minimize disk I/O, KVSWARE estimates which KV entries are most important for attention computing in the next layers and only

⁴From 7.6s to 37.2s on Llama3-8B with NVIDIA Orin AGX, batch size = 1.

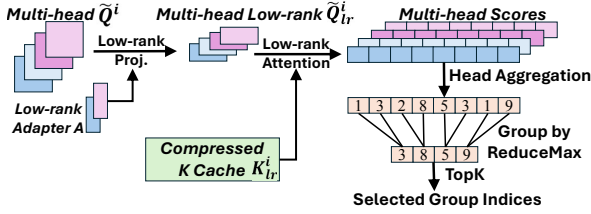


Figure 6: Grouped critical KV entries prediction.

preloads these. A unique feature of KVSWAP is that it first groups G consecutive key-value (KV) entries to align with the block-read characteristics of devices like NVMe, UFS, and eMMC (see also Fig. 2), then predicts and loads the most important groups. This is different from prior work [41, 59, 60] that predicts on individual heads or tokens. The group size is configurable and can be automatically determined during offline processing (Sec. 3.5).

Fig. 6 depicts how we predict which *groups* of KV cache entries to be loaded from disk. The idea is to use the low-rank K cache, K_{lr} , to estimate the attention scores (QK^T) without accessing the full K cache. K cache elements (and their associated values, V) that contribute to a high attention score will be prefetched into memory. Specifically, for query head h , we approximate the attention score as:

$$Q_h K_{g(h)}^T \approx Q_h (K_{lr} A_{g(h)}^T)^T = (Q_h A_{g(h)}) K_{lr}^T \quad (1)$$

Here, $Q_h \in \mathbb{R}^{1 \times d}$ is the query vector for head h ; $A_{g(h)} \in \mathbb{R}^{d \times r}$ is the low-rank adapter for the K cache assigned to h ; and $g(h)$ maps each query head to its shared K head. We refer to the term $Q_h A_{g(h)}$ as a low-rank query vector. This enables the reconstruction of multi-head attention from a compressed, head-unified K_{lr} , thereby reducing computation and allowing us to decide which *groups* of KV entries are worth loading.

Online prediction. We predict the critical KV groups ahead of time to issue disk load concurrently with the computation. Specifically, while computing layer $i-1$, KVSWAP predicts in advance the critical KV groups needed for layer i . To do this, we approximate the upcoming input X_i using X_{i-1} , taking advantage of the observation that the input (embeddings) to a Transformer layer i is often similar to that at layer $i-1$ [41, 65]. Using this approximation, we apply layer i 's projections to obtain a low-rank multi-head query vector $\tilde{Q}_{lr}^i$. Together with K_{lr}^i , this vector is used in a low-rank attention computation to estimate the attention scores for layer i . Since early-layer inputs vary significantly, we use exact inputs for the first two layers and previous-layer inputs thereafter, as justified in Sec. 6.1.

Scoring and selection. To select critical KV cache entries, we first compute a single importance score for each token by summing the scores across all heads. We then form groups of G consecutive tokens and represent each group by the maximum score among its members. Finally, we select the top-ranked groups with the highest representative scores. A ReduceMax operation within each group captures the most salient token score, and a TopK step selects the indices of M most relevant KV entry groups to load for layer i .

Compared with prior prediction method. Unlike prior work [41], our approach introduces prediction at the group granularity, rather than at individual KV entries. Beyond this, while InfiniGen [41], a representative prior work, also employs approximated attention

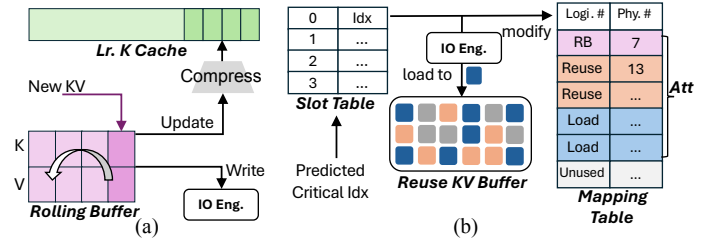


Figure 7: KVSWAP runtime: (a) a compressed K cache and (b) reuse buffer for KV cache management.

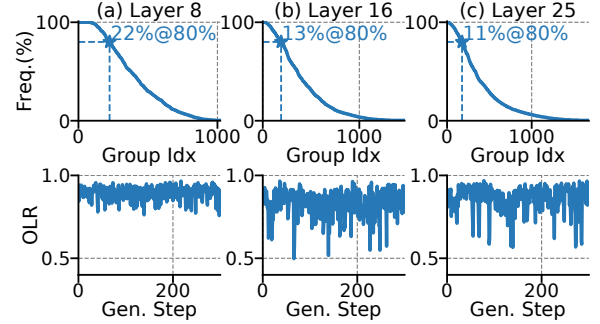


Figure 8: Frequency and overlap ratio (OLR) of predicted critical KV groups over 300 decoding steps. Group indices are reordered by frequency for clarity.

as its prediction mechanism, our algorithm differs in two fundamental aspects. (1) InfiniGen adopts an index-selecting strategy that pre-determines which embedding dimensions of the K cache to retain and directly selects them via indices. Under stringent K cache compression, this strategy struggles to preserve fidelity as the number of selected indices decreases. (2) InfiniGen applies multi-head attention directly to the index-selected K cache, whereas our method constructs a head-unified low-rank K cache that aggregates information across heads, enabling more aggressive compression. We then reconstruct multi-head attention from this representation using Eq. 1.

3.4 KVSWAP Runtime System

We now introduce the overall KVSWAP runtime system, which operates during the prefilling and decoding stages (Sec. 2.1). During the *prefilling* phase, KVSWAP captures the LM-generated KV cache for the input prompts and writes it to disk in a layer-by-layer fashion. At the same time, it computes a low-rank K cache from the full K cache to create an initial compact K cache (Sec. 3.2) that is stored in memory. When *decoding* begins, the runtime coordinates the grouped critical KV predictor (Sec. 3.3) and the KV cache manager (Sec. 3.4.4) to efficiently support disk-based KV cache loading and attention computation. It identifies, loads, and reuses important KV entry groups while keeping the most recent KV entries in memory. To maximize efficiency and overlap I/O with computation, these operations for the next layer are performed concurrently with the attention and feed-forward network (FFN) computations of the current layer (Sec. 3.3).

3.4.1 New KV entries management. As shown in Fig. 7a, our in-memory KV cache has two components: a low-rank compressed

K cache (Sec. 3.2) and a rolling buffer (RB). The compressed K cache serves as the input to our KV predictor (Sec. 3.3). During decoding, new entries are appended to the RB. Since critical entries are predicted in groups, the importance of recent entries cannot be assessed until a group is completed. Thus, the RB temporarily stores recently generated KV entries until they accumulate to a full group of size G for offloading to disk.

3.4.2 Locality in grouped KV predictions. We observe that a subset of critical KV entries exhibits temporal locality across layers. For example, applying our grouped KV predictor (with group size=4) to Qwen3-8B on QMSum [15] samples, we define the *overlap ratio* at generation step j as the fraction of important groups at step j that also appeared at step $j-1$. Fig. 8 reports frequency and overlap ratios over 300 decoding steps⁵. Fewer than 22% of groups account for 80% of occurrences, showing that the set of important groups changes substantially over time. However, adjacent steps exhibit strong overlap, revealing temporal redundancy. This motivates our reuse strategy: retaining recently accessed critical groups in memory allows them to be efficiently reused in subsequent steps without reloading from disk.

3.4.3 Reuse buffer for caching grouped KV entries. We introduce a *KV reuse buffer* (not to be confused with the rolling buffer described in Sec. 3.4.1 that stores the recently generated KV entries) to store recently accessed critical KV entries (loaded from disk) for potential reuse across prediction steps. By reusing overlapping groups of critical KV entries rather than reloading them, the reuse buffer reduces I/O overhead - especially on low-bandwidth disks or when handling large KV volumes. We implement the reuse buffer as a fixed set of dedicated *memory slots*, each capable of storing one KV group. A slot table records the group ID currently stored in each slot. When the predictor identifies the next set of critical groups (Fig. 7b), the system first checks the slot table to see whether a group is already in the reuse buffer. If present, the corresponding slot from the reuse buffer is passed to the attention kernel for computation directly; if not, the KV group is loaded from disk and stored in an available slot, with the slot table updated accordingly. We use a simple FIFO policy as the replacement strategy for the reuse buffer.

3.4.4 KV cache manager. With KVSWARE, attention computation consumes KV entries from: hit slots in the reuse buffer, important KV groups that are not in the reuse buffer and need to be loaded from disk, and recent newly generated entries from the rolling buffer. The KVSWARE KV cache manager uses a mapping table to logically address these heterogeneous memory regions - similar to OS virtual memory. This abstraction provides direct compatibility with PagedAttention [39], which expects a contiguous logical view of the KV cache. During generation, reuse patterns change over time, causing the valid memory regions and the positions of usable data in the reuse buffer to shift. To handle this, the runtime updates the mapping table before each attention computation to reflect the current KV entry layout (Fig. 7b).

⁵Results from layers 8, 16, 25 are shown; other layers have similar patterns.

3.5 Offline Parameter Tuning

KVSWARE provides an API (Fig. 4a) for *one-off* offline parameter tuning, which selects optimal runtime settings based on user constraints (maximum memory $\mathcal{B}_{\max}$, context length $S_{\max}$, and batch size $b_{\max}$), the target LM, and the hardware platform. In our evaluation, tuning completes in under 30 minutes and outputs a JSON file that records the best-found parameters for runtime scheduling. These include group size for KV prediction G , K-cache compression ratio σ , number of selected groups M , and reuse buffer capacity C .

Parameter search tries to balance three tightly coupled factors: *generation quality*, *inference throughput*, and *memory usage*. For example, G , M , and C affect both I/O cost and computation, directly influencing throughput. G and σ shape generation quality, while σ , M , and C determine memory footprint. These interdependencies create a three-way trade-off where improving one dimension often degrades another. Since accurate quality evaluation is also costly, finding a global optimum is impractical. To address this, we use a *greedy-based parameter solver* to efficiently search offline for good local configurations. More details of our parameter search method can be found in Appendix A.

3.6 Implementation Details

We implement the core component of KVSWARE by extending FlexGen [58] with about 5K lines of Python code. LM computation uses the PagedAttention kernel from vLLM [39] together with PyTorch [54] operators. The KVSWARE predictor and main computation execute asynchronously, with synchronization mechanisms to ensure correctness.

For disk access, we develop the underlying I/O engine using the asynchronous I/O interface *io_uring* provided by the Linux kernel, with a small CPU-pinned memory buffer as a staging area between disk and GPU memory. At initialization, KV cache files are preallocated on disk according to a user-configured size. The KV cache is stored on disk in a grouped layout with a configurable block size G , where the keys and values of all heads in a group are placed contiguously to improve locality and effective disk bandwidth.

For KV reuse, we maintain a mapping table that marks valid KV positions. Missed entries are transferred in bulk to contiguous GPU memory rather than scattered back to their original slots, avoiding many small memory copies and reducing data movement overhead.

4 Experimental Setup

4.1 Platform and Workloads

Evaluation platform. Our main evaluation platform is the NVIDIA Jetson Orin AGX [1] embedded platform with a 12-core ARM Cortex-A78 CPU, an Ampere GPU with 2048 CUDA cores and 64 GB of unified RAM. In addition, we also evaluate KVSWARE on the NVIDIA Jetson Orin Nano, which features a reduced 6-core CPU, 1024 CUDA cores, and 8 GB of RAM, to demonstrate cross-platform generalization (Sec. 5.2.2).

We apply KVSWARE to two widely used disk types on mobile and embedded systems: NVMe (I/O bandwidth 1.8 GB/s) and eMMC (I/O bandwidth 250 MB/s). The system runs JetPack-6.2 with Linux kernel 5.15.148-tegra, CUDA 12.6 and PyTorch 2.7. KVSWARE is designed for scenarios with a tight memory budget. However, because

some competing baselines require substantial memory, we leverage the platform’s 64 GB of RAM to ensure all methods are evaluated fairly. Sec. 4.3 reports the actual KV-cache memory used in each experimental setting.

Language models. We evaluate KVSWAP on both text and video understanding tasks using eight LMs of varying sizes. For text tasks, we use Llama-3.1-8B-Instruct (Llama3-8B), Llama-3.2-3B-Instruct (Llama3-3B) [22], and Qwen3 [74] models (0.6B-14B). We enable the thinking mode of Qwen3 to evaluate it on CoT reasoning scenarios. For video understanding, we select Qwen2.5-VL-3B, 7B [14] and InternVL3-14B [86].

Tasks. We test on long-context tasks using all English tasks from LongBench [15], Needle-in-a-Haystack (NIAH) [37], and eight tasks from RULER [31]. Test context length is fixed at 32K for RULER, and ranges from 4K to 32K for LongBench and NIAH. For video understanding, we use open-ended generation tasks from MLVU [85], including Sub-Scene Captioning (SSC) and Video Summarization (VS), with context lengths of 22K–32K per video.

4.2 Competing Baselines

We compare KVSWAP against an extensive set of baselines and prior work, as described below.

InfiniGen [41]: This closely related work offloads KV cache between GPU and CPU memory by predicting and selecting important KV entries *per* head and token position, unlike our group-based prediction. We extend InfiniGen to the modern GQA attention scheme and adapt its runtime for disk-based offloading. We also evaluate three enhanced variants of InfiniGen: 1) *InfiniGen** incorporates head aggregation from our prediction strategy (Fig. 6), 2) *InfiniGen*+ru* further introduces our KV reuse strategy to reduce disk I/O overhead, and 3) *InfiniGen*+gp* additionally applies our grouped selection mechanism.

FlexGen [58]: It is designed to offload the entire KV cache to disk. During decoding, the full KV cache is restored layer by layer into memory for full attention computation.

ShadowKV [60]: This CPU-GPU KV offloading framework stores a low-rank K cache on the GPU and offloads the V cache to CPU memory (see also Sec. 3.2). During generation, it loads only important parts of the V cache and reconstructs the K cache on-the-fly. Like InfiniGen, we also adapt its runtime to support disk offloading.

Loki [59]: It is a sparse attention method originally designed to accelerate attention computation. We modify its core approximate attention formulation to serve as a predictor for identifying critical KV entries and embed it into our modified version of InfiniGen with disk-offloading support, replacing its core prediction components.

vLLM [39]: This industry-grade LM serving system stores the full KV cache in memory by default. We use vLLM to *approximate the throughput upper bound*, measuring how closely an offloading strategy approaches the ideal case without disk overhead. All remaining device memory - after storing model weights and allocating workspace - is dedicated to the KV cache. We tune vLLM parameters and report the best throughput achieved.

SnapKV [42]: In Sec. 6.3, we compare KVSWAP to this representative in-memory compression method that compresses the prompt KV entries to save memory.

Table 1: Tested KV memory budgets for Llama3-8B.

Settings	KV Memory Configuration (MiB)		vLLM (GiB)
	Relaxed	Tight	
A (Sec. 5.1, 5.2, 5.3)	310/batch@32K	120/batch@32K	31
B (Sec. 5.4)	2000 in total	800 in total	

Note: For KVSWAP in both settings (and offloading baselines in Setting A), we constrain the *per-batch* KV budget of all evaluated LMs to 1/13 (relaxed) and 1/34 (tight) of their full KV cache. In Setting B, we fix the *total* KV budget to 2000 MiB (relaxed) and 800 MiB (tight) for all offloading methods.

4.3 Evaluation Methodology

Perspective. We assess KVSWAP under two settings. **(A)** We align the *per-batch* memory consumption of each method to a consistent, low budget and compare their performance with varying conditions under this constraint. **(B)** We fix the *total* memory budget and evaluate each method under its best-case configuration. Notably, for baselines at their optimal settings, per-batch memory consumption varies, which directly impacts the maximum achievable batch size. We evaluate setting A in Sec. 5.1, 5.2, and 5.3, and setting B in Sec. 5.4.

Setting A. In setting A, for all offloading methods, we constrain the *per-batch* KV memory budget of all evaluated LMs to fixed fractions of their corresponding full KV cache. We employ two budgets: a relaxed budget (1/13 of the full KV cache) and a tight budget (1/34 of the full KV cache). The latter is denoted with the suffix “-t” (e.g., KVSWAP-t). Specifically, we reconfigure the baselines as follows: InfiniGen by adjusting the partial weight ratio, Loki by tuning key dimensionality, and ShadowKV by adjusting chunk size, outliers, and K-cache rank. For KVSWAP, we set $S_{\max} = 32K$, $b_{\max} = 16$, $MG = 400$, and maximum K-cache compression ratio $\sigma_{\max} = 32$.

Setting B. In setting B, we apply the same configuration for KVSWAP as in setting A. And for all offloading baselines, we apply their optimal configurations. Specifically, for ShadowKV and LoKi, we directly adopt the configurations reported in their source publications. For InfiniGen, since no reference configurations are readily available, we experimentally set the partial weight ratio to 0.5, resulting in a conservative 4× KV cache memory reduction. Here, their per-batch memory consumption varies and is generally higher than that of KVSWAP. To ensure a fair comparison, we instead fix the *total* memory budget and evaluate each method at the *maximum batch size* it can support within this budget. We evaluate under a relaxed total budget of 2000 MiB and a tight total budget of 800 MiB.

Example test memory configuration. Tab. 1 shows the tested memory budgets for Llama3-8B. For vLLM, it is an ideal baseline that uses all available device memory for KV cache (31 GiB) and is not affected by the evaluation settings. Qwen3-8B has a memory budget similar to that of Llama3-8B.

4.4 Performance Metrics

We report *generation accuracy* and *throughput* (tokens per second) during decoding. Accuracy is reported under varying KV cache budgets and context lengths. Throughput is measured as the average decoding rate over 1000 continuously generated tokens. Each result is the average of five runs with different randomly selected test inputs.

5 Experimental Results

5.1 LM Generation Quality

5.1.1 Text processing. Tab. 2 reports the generation accuracy of Llama3-8B on RULER and LongBench with maximum context length 32K. *Full-KV* shows raw accuracy using the entire KV cache, while other results indicate accuracy loss relative to Full-KV. KVSWAP occasionally exceeds Full-KV due to the probabilistic nature of LMs. Results for Qwen3-8B are similar and given in Tab.1 of Appendix B.

KVSWAP and KVSWAP-t (with a tighter memory budget - see Sec. 4.3) outperform all offloading baselines with acceptable accuracy loss. InfiniGen has the largest degradation, as it fails to capture critical KV entries under tight budgets. Although InfiniGen* reduces accuracy loss by incorporating our head aggregation to reduce prediction noise, it still falls far behind KVSWAP, with average accuracy losses of 78.7% on RULER and 14.2% on LongBench. InfiniGen*+gp, which further applies our grouped selection strategy, still suffers from similarly substantial accuracy degradation. By contrast, KVSWAP keeps average accuracy loss within 4.4% and 1.1%. ShadowKV leads to 52.3% and 5.0% accuracy loss. Loki also suffers high losses of 34.0% and 8.8%. Furthermore, under a highly constrained budget, only KVSWAP-t maintains usable accuracy: $\leq 5.6\%$ and 2.4% loss for NVMe. Note that KVSWAP^{NVMe} and KVSWAP^{eMMC} differ because the best group size is $G=4$ for NVMe and $G=8$ for eMMC.

Fig. 9 compares KVSWAP-t with Loki-t and ShadowKV-t on NIAH with Qwen3-8B under a tight 130 MiB per-batch memory budget (using NVMe). The x-axis shows the context length (or token limit) to be tested, and the y-axis shows the relative position to be tested within a context; lower scores with dark regions indicate areas where the model exhibits degraded capability. Only KVSWAP-t maintains full processing capability at all sequence locations, while Loki-t and ShadowKV-t suffer substantial generation performance degradation under the same memory budget.

5.1.2 Reasoning and video understanding. Tab. 3 reports generation quality on CoT reasoning and video understanding LMs. For reasoning, we evaluate on three multi-document QA tasks from LongBench (HotpotQA, 2WikiMultihopQA, MuSiQue) [15], reporting averaged accuracy across tasks. For video understanding, we evaluate Qwen2.5-VL-3B (QVL3B), 7B (QVL7B), and InternVL3-14B (IVL14B), reporting averaged scores on MLVU-SSC and MLVU-VS (range: 2-10). Recalling evaluation setting A (Sec. 4.3), we configure all offloading methods with per-batch memory budgets set to 1/13 and 1/34 of the full KV cache, applied across all models. We omit InfiniGen and InfiniGen* since Tab. 2 and Tab. 1 of Appendix B already confirm their poor performance.

Like with text processing, KVSWAP outperforms Loki and ShadowKV on NVMe and eMMC disks, with minor accuracy losses $\leq 4.6\%$, and score losses ≤ 0.4 . KVSWAP-t also strikes a good balance between KV cache memory footprint and generation quality, while all other methods incur severe degradation in output quality ($\geq 45.0\%$ accuracy loss on reasoning models and ≥ 2.1 points, or 46.7% , on video models). In contrast, KVSWAP-t constraints losses to 3.6-10.0% on reasoning models and 0.4-1.8 points on video models. Its advantage grows with model size: ≥ 0.3 points on QVL3B, ≥ 1.7 on QVL7B, and ≥ 1.9 on IVL14B.

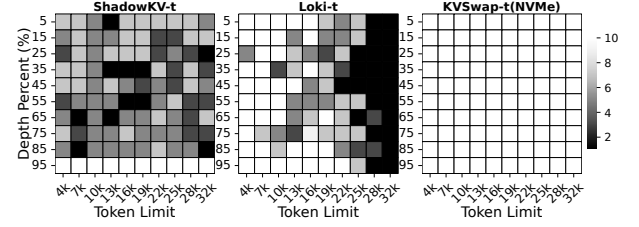


Figure 9: Qwen3-8B generation quality on NIAH. x-axis is the tested context length and the y-axis is the relative position to test within a context length.

Summary. KVSWAP maintains accuracy within an acceptable range across tasks, model sizes, modalities, context lengths, and disk types, consistently outperforming KV cache offloading baselines under the same low-level memory budget. Moreover, KVSWAP-t remains robust under a very tight memory budget, while most baselines fail and become impractical for use. These results highlight an important insight: designs like KVSWAP that aggressively reduce memory while preserving prediction accuracy are crucial for accurate long-context inference under limited memory budgets.

5.2 LM Generation Throughput

5.2.1 Throughput on Orin AGX. Tab. 4 presents generation throughput of Llama3-8B on Orin AGX with batch sizes ranging from 1 to 16 and context lengths of 16K and 32K. All, except FlexGen and vLLM, use a fixed per-batch KV memory budget of 1/13 of the full KV cache (see setting A in Tab. 1). FlexGen is unconstrained, as it offloads the full KV cache, while vLLM serves a throughput baseline under idealized memory. Loki and InfiniGen perform similarly due to their fine-grained per-token and per-head design, which yields poor I/O efficiency; their results are averaged for brevity. KVSWAP outperforms all other offloading methods across storage devices while maintaining the best generation quality (Sec. 5.1). Throughput remains stable across context lengths from 16K to 32K since the number of selected KV entries is fixed, making I/O latency largely independent of sequence length.

FlexGen, Loki, and InfiniGen give only up to 1.9 and 0.1 tokens/s on NVMe and eMMC, respectively, due to heavy disk traffic (FlexGen) or poor bandwidth utilization (Loki and InfiniGen). KVSWAP sustains 6.9 tokens/s (NVMe) and 5.9 tokens/s (eMMC) with a batch size of 1 at 16K context, rising to 35.1 and 15.8 tokens/s at a batch size of 8. On NVMe, throughput continues to scale, reaching 46.1 tokens/s at a batch size of 16, whereas eMMC saturates due to its limited bandwidth. Even under saturation, KVSWAP maintains significant speedups over competing baselines. InfiniGen* and InfiniGen*+ru provide moderate improvements but remain significantly behind KVSWAP. Although InfiniGen*+ru+gp achieves notable speedups on eMMC and approaches KVSWAP, it still lags significantly on NVMe. At batch sizes of 16 and a 32K context, it is $2.5\times$ slower than KVSWAP, with substantially lower accuracy (Tab. 2).

Finally, at 32K context with batch size 1, KVSWAP uses $13\times$ less memory while achieving 75% and 65% of vLLM's throughput on NVMe and eMMC, respectively. At larger batch sizes, KVSWAP can even outperform vLLM on NVMe by carefully trading generation accuracy for throughput. For example, at 16K context and batch size 16, it is $1.2\times$ faster while using $11.0\times$ less KV memory. Whereas

Table 2: Relative accuracy loss (%) of Llama3-8B compared to Full-KV on RULER (left) and LongBench (right).

Methods	S1	S2	MK1	MQ	MV	QA1	QA2	VT	Avg.	SQA	MQA	SUM	FSL	SYN	COD	Avg.
Full-KV (raw %)	100.0	100.0	100.0	98.8	99.5	82.0	56.0	96.6	91.6	39.3	41.4	28.0	68.0	99.5	49.1	54.2
InfiniGen	-100.0	-100.0	-100.0	-98.8	-99.5	-61.0	-40.0	-96.6	-87.0	-36.0	-37.8	-26.8	-50.8	-99.0	-44.7	-49.2
InfiniGen*	-86.0	-97.0	-97.0	-98.3	-99.5	-39.0	-20.0	-93.0	-78.7	-14.4	-14.8	-10.7	-14.6	-5.3	-25.4	-14.2
InfiniGen*+gp	-53.0	-98.0	-97.0	-97.6	-99.3	-34.0	-10.0	-76.4	-70.1	-13.5	-11.2	-14.2	-15.9	-7.5	-22.4	-14.1
Loki	-3.0	-14.0	-18.0	-67.5	-61.5	-28.0	-15.0	-65.4	-34.0	-7.4	-4.9	-10.8	-8.0	-0.3	-21.6	-8.8
ShadowKV	0.0	-86.0	-84.0	-93.0	-94.5	-13.0	-7.0	-41.0	-52.3	-9.7	-2.8	-4.1	-5.1	-1.5	-6.8	-5.0
KVSWAP ^{NVMe}	0.0	-1.0	0.0	-3.3	-7.0	-1.0	+2.0	-10.4	-2.6	+0.9	-0.3	-2.1	-0.7	0.0	-1.8	-0.6
KVSWAP ^{eMMC}	0.0	0.0	0.0	-9.0	-12.0	0.0	-3.0	-11.8	-4.4	-0.5	-0.1	-2.3	-1.4	-0.5	-1.8	-1.1
Loki-t	-99.0	-100.0	-100.0	-98.8	-99.5	-57.0	-32.0	-96.6	-85.4	-24.5	-25.7	-20.5	-43.2	-7.6	-38.2	-26.6
ShadowKV-t	-21.0	-89.0	-93.0	-95.8	-95.0	-16.0	-6.0	-79.2	-61.9	-15.5	-7.0	-7.1	-10.2	-4.5	-8.0	-8.7
KVSWAP-t ^{NVMe}	0.0	-1.0	-3.0	-10.8	-13.3	-4.0	0.0	-13.0	-5.6	-2.5	-0.4	-3.0	-2.7	0.0	-5.8	-2.4
KVSWAP-t ^{eMMC}	0.0	-13.0	-16.0	-54.3	-55.8	-13.0	-7.0	-24.6	-22.9	-4.2	-1.1	-3.9	-1.5	-0.5	-5.2	-2.7

Table 3: Relative accuracy (%) and score loss over Full-KV on reasoning (left) and video understanding (right).

Methods	Q4B	Q8B	Q14B	QVL3B	QVL7B	IVL14B
Full-KV (raw %)	60.5	62.5	66.0	4.5	4.9	4.8
Loki	-10.0	-4.1	-9.6	-1.4	-1.0	-2.0
ShadowKV	-6.6	-5.4	-7.3	-0.7	-0.6	-0.5
KVSWAP ^{NVMe}	-1.8	-2.4	-2.6	-0.4	-0.3	-0.3
KVSWAP ^{eMMC}	-4.0	-3.0	-4.6	-0.4	-0.3	-0.3
Loki-t	-57.2	-47.2	-58.6	-2.5	-2.9	-2.7
ShadowKV-t	-50.1	-47.5	-45.0	-2.1	-2.4	-2.4
KVSWAP-t ^{NVMe}	-3.6	-4.8	-6.7	-1.7	-0.6	-0.4
KVSWAP-t ^{eMMC}	-7.2	-6.9	-10.0	-1.8	-0.7	-0.5

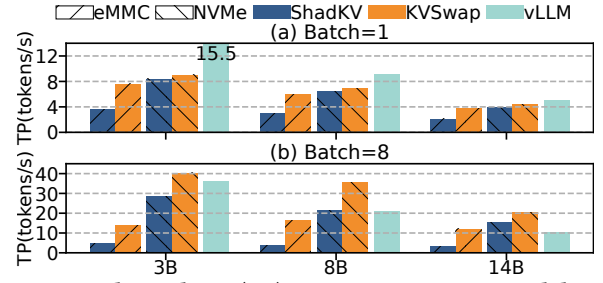
Table 4: Throughput (tokens/s) comparison of Llama3-8B on Orin AGX across batch sizes and context lengths (CLs).

Methods	CL=16K					CL=32K				
	b=1	b=2	b=4	b=8	b=16	b=1	b=2	b=4	b=8	b=16
eMMC	FlexGen	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
	Loki/InfiGen	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
	InfiGen*	1.7	1.7	1.5	1.4	1.6	1.5	1.4	1.4	1.4
	InfiGen*+ru	4.0	4.7	5.2	4.2	3.7	4.2	4.4	4.9	4.5
	InfiGen*+ru+gp	5.3	7.9	12.9	12.6	10.5	5.0	7.7	13.0	12.1
	ShadowKV	3.0	3.8	4.1	4.4	3.4	3.0	3.7	4.2	3.8
NVMe	KVSWAP	5.9	10.5	16.2	15.8	11.2	6.0	10.3	15.2	15.7
	FlexGen	0.8	0.8	0.8	0.8	0.8	0.4	0.4	0.4	0.4
	Loki/InfiGen	1.9	1.9	1.9	1.9	1.9	1.9	1.9	1.8	1.8
	InfiGen*	5.0	8.4	10.8	11.3	13.1	5.1	7.5	8.9	10.7
	InfiGen*+ru	5.2	9.2	14.9	22.5	26.6	4.9	8.5	11.9	14.3
	InfiGen*+ru+gp	5.8	10.2	16.2	25.3	27.4	5.1	8.5	12.6	15.4
vLLM	ShadowKV	6.4	10.4	16.3	21.9	26.7	6.4	10.0	16.3	21.5
	KVSWAP	6.9	11.9	20.8	35.1	46.1	6.9	11.9	21.4	35.6

Table 5: Throughput (tokens/s) comparison of Qwen3-0.6B and Qwen3-1.7B on Orin Nano across batch sizes at 16K context length.

	Methods	Qwen3-0.6B				Qwen3-1.7B			
		b=1	b=2	b=4	b=8	b=1	b=2	b=4	b=8
NVMe	InfiGen*+ru+gp	7.6	14.2	18.8	22.8	7.3	12.8	18.9	22.9
	ShadowKV	8.1	12.1	16.0	18.5	6.6	11.9	15.9	17.7
	KVSWAP	7.6	14.8	24.6	36.7	7.3	14.6	23.0	33.5
-	vLLM	37.4	28.9	30.3	31.7	20.0	17.0	17.2	17.6

vLLM saturates once its 31 GiB cache limit is exceeded, KVSWAP delivers up to 2.0× higher throughput with substantially less KV cache memory at an increased 32K context length, a benefit that grows with workload scale. Additional comparisons for 8K and 24K context lengths are provided in Tab. 2 of Appendix B.

**Figure 10: Throughput (TP) comparison across model sizes at a 32K context with batch sizes of (a) 1 and (b) 8.**

5.2.2 Throughput on Orin Nano. Tab. 5 reports generation throughput on Orin Nano with NVMe for Qwen3-0.6B and Qwen3-1.7B, with batch sizes from 1 to 8 and a 16K context length. Consistent with Orin AGX, KVSWAP generalizes well across platforms and outperforms offloading baselines in most cases.

For Qwen3-0.6B at batch size 8, KVSWAP achieves 36.7 tokens/s, outperforming InfiniGen*+ru+gp, ShadowKV, and vLLM by 1.6×, 2.0×, and 1.2×, respectively. For Qwen3-1.7B, KVSWAP further delivers a 1.9× speedup over vLLM. Notably, KVSWAP achieves similar throughput on both models because their KV cache offloading sizes are identical, and I/O dominates the execution time.

5.2.3 Summary. KVSWAP significantly improves throughput over existing KV cache offloading methods across context lengths, platforms, and storage types. These results demonstrate the effectiveness of leveraging grouped I/O access and KV reuse to enable high-throughput KV cache disk offloading. Furthermore, KVSWAP outperforms vLLM under heavy workloads while using substantially less KV memory, highlighting the benefits of our aggressive KV memory compression for memory-efficient execution.

5.3 Impact of LM Model Sizes

Fig. 10 compares throughput on Orin AGX for Llama3 (3B and 8B) and Qwen3-14B at a 32K context length with batch sizes of 1 and 8, given available KV cache memory for vLLM of around 43, 31, and 16 GiB. Across model sizes, KVSWAP outperforms ShadowKV on eMMC, achieving speedups of over 1.8× (up to 2.1×) at batch size 1, and over 2.9× (up to 4.1×) at batch size 8. On NVMe, KVSWAP achieves speedups of at least 1.3× (up to 1.8×) at batch size 8. Although throughput is comparable to ShadowKV on NVMe with batch size 1, KVSWAP delivers much better generation quality, as

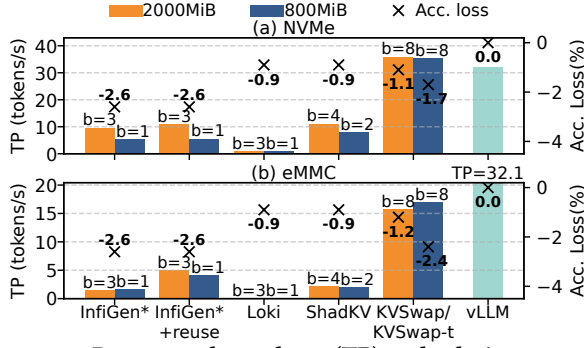


Figure 11: Best-case throughput (TP) and relative accuracy loss for Llama3-8B using (a) NVMe and (b) eMMC. b = largest possible batch size.

discussed in Sec. 5.1. When compared with vLLM at batch size 8, KVSWARE (NVMe) achieves throughput improvements of 1.1 $\times$, 1.7 $\times$, and 1.9 $\times$ on the 3B, 8B, and 14B models, respectively. Notably, on the 14B model, KVSWARE with the low-end eMMC disk even surpasses vLLM by 1.2 $\times$ at batch size 8. These results demonstrate that KVSWARE scales well as the model grows larger.

5.4 Best-case Configurations for Baselines

So far, we have set a consistent per-batch memory budget for each KV cache offloading scheme. We now compare KVSWARE against offloading baselines using their recommended best-case configurations. Given total memory budgets of 2000 MiB and 800 MiB, we compare throughput using the largest possible batch size supported by each method, as described in Sec. 4.3-setting B. We also report the average accuracy loss relative to Full-KV on two LongBench task types, multi-document QA (MQA) and summarization (SUM), covering six tasks in total. For reference, Full-KV achieves an average accuracy of 34.7% on these challenging tasks for Llama3-8B.

Fig. 11 presents the results on Orin AGX, with vLLM numbers taken from Sec. 5.2 for direct comparison. KVSWARE proves highly effective at trading a small loss in accuracy for substantially higher throughput and much lower memory usage. While vLLM follows the Full-KV design without accuracy loss, ShadowKV and Loki achieve the next two lowest accuracy losses, they do so at the cost of either high memory demand or low throughput. In contrast, KVSWARE delivers clear advantages in both throughput and memory efficiency, with only marginal accuracy degradation. Compared with vLLM, KVSWARE on NVMe reduces KV cache memory by 12.3 $\times$ –30.7 $\times$ while still providing 1.1 $\times$ higher throughput, at a maximum accuracy drop of just 2.4%. Relative to ShadowKV, KVSWARE achieves large throughput gains - 7.1 $\times$ and 8.6 $\times$ on eMMC, and 3.3 $\times$ and 4.5 $\times$ on NVMe - with accuracy drops $\leq 1.5\%$.

5.5 Larger Memory Budget

In this experiment, we allocate a 1 GiB *per-batch* KV memory budget to evaluate if KVSWARE can support long context with GiB-level memory budgets. Tab. 6 shows the accuracy of several tasks on RULER with a 128K context using Llama3-8B.

For InfGen* and ShadowKV, we follow the best-case configurations described in Sec. 5.4. As shown in Tab. 6, only KVSWARE supports the full 128K context under this memory budget, while

Table 6: Accuracy (%) and supported context length (CL) on RULER (128K) for Llama3-8B under a 1 GiB per-batch KV budget.

Subtask	InfGen*	ShadowKV	KVSWARE	Full-KV
S1	9.0	44.0	87.0	100.0
S2	0.0	46.0	87.0	100.0
MK1	1.0	42.0	85.0	97.0
Supported CL	32K	55K	128K	-

Table 7: Accuracy (%) comparison for cross-domain robustness of KVSWARE on Qwen3-8B.

Task	C4-profiled	In-domain-profiled	Full-KV
MATH-500	92.0	92.2	92.4
LiveCodeBench-v5	55.7	54.9	56.9

other methods are limited to much shorter contexts. For example, ShadowKV supports only 55K. This translates to substantial accuracy differences. KVSWARE remains closest to Full-KV, maintaining over 85% accuracy, whereas InfGen* nearly collapses (≈ 0) and ShadowKV achieves less than 50%.

5.6 Domain Robustness

We build our low-rank adapter (Sec. 3.2) using samples from a general-domain dataset (e.g., C4 [57]). To evaluate its cross-domain robustness, we apply the C4-profiled adapter to MATH-500 [43] and LiveCodeBench-v5 [33], and compare it with Full-KV and adapters profiled on the corresponding domains. Tab. 7 reports the accuracy on Qwen3-8B (thinking). As shown, C4-profiled KVSWARE matches in-domain profiling within 0.8% and incurs $\leq 1.2\%$ accuracy loss compared to Full-KV across both tasks, demonstrating strong cross-domain robustness.

5.7 Multi-programmed Workloads

We now evaluate KVSWARE in a multi-programmed setting alongside two common mobile workloads: video playback and object detection. For video playback, we use CPU-based software (i.e., ffmpeg) decoding of 1080p streams. For object detection, we run YOLO11x, the largest Ultralytics YOLO11 model [8], processing 640 $\times$ 640 inputs from a 30 FPS USB camera. To stress resource contention, we avoid using the on-chip NVDEC/VIC units for video decoding and the DLA for YOLO execution. Throughput is measured with Llama3-8B at batch size 8 and 16K context length on Orin AGX.

Video playback reduces KVSWARE's throughput by only 17.2% and 15.0% on NVMe and eMMC, respectively, suggesting that it can efficiently co-run with UI workloads. For object detection, the throughput drops by 48.1% on NVMe and 28.2% on eMMC. This is expected as the competing workload increases. Nevertheless, on eMMC, KVSWARE still maintains a relatively stable performance. On NVMe, we attribute the larger drop to more severe GPU compute contention, which also causes a 47.2% drop for vLLM.

6 System Analysis

6.1 Design Choices

To evaluate how our key design choices, including group size, KV reuse buffer, number of selected KV entries, rolling buffer size and layer-to-layer approximation, affect the overall throughput and

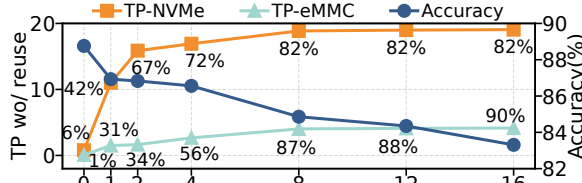


Figure 12: Trade-off between accuracy, I/O efficiency, and throughput (TP) under varying group sizes (x-axis).

Table 8: Reuse ratio and throughput (TP) statistics. "No Reu." = TP wo/ reuse. "↑" = TP improvement of w/ over wo/ reuse.

Disk	Dataset	Metric	Min	Max	Std	Avg	No Reu.	↑
NVMe	QMSum	Reuse(%)	76.2	79.1	0.7	77.3	-	2.1×
		TP(toks/s)	33.4	38.7	2.0	35.8	17.3	
	MSQue	TP(toks/s)	32.4	37.6	1.5	35.4	17.4	2.0×
eMMC	QMSum	Reuse(%)	76.1	79.2	0.7	77.2	-	4.0×
		TP(toks/s)	14.4	17.5	0.6	15.7	3.9	
	MSQue	TP(toks/s)	14.5	16.0	0.4	15.3	4.0	3.8×

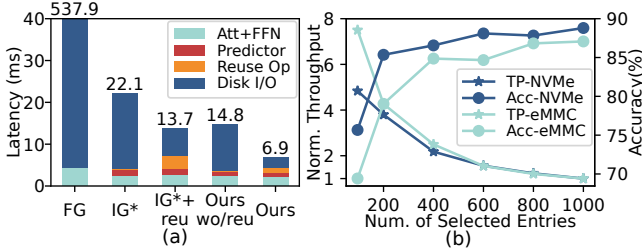


Figure 13: (a) Decoding latency breakdown on NVMe. (b) Trade-off between accuracy and throughput across KV selection sizes and disk types.

accuracy, we apply KVSWAP to Llama3-8B under a 310 MiB per-batch KV memory budget. Accuracy is reported as the average over MV, QA1, and VT tasks from RULER, while throughput and latency are measured at a 32K context length with a batch size of 8.

Group size and system efficiency. Fig. 12 shows how throughput and accuracy vary with KV prediction group sizes (G). The number beside each point indicates I/O utilization. A group size of 0 disables our head aggregation strategy. This ablation study reports throughput *without* KV reuse to isolate our optimizations. As G increases, accuracy drops gradually from 88.8% to 83.3%, while throughput (without reuse) improves from 1.8 to 19.1 tokens/s on NVMe and 0.1 to 4.2 tokens/s on eMMC. When $G = 0$ or 1, both throughput and I/O utilization are low. This highlights that our grouped critical KV prediction is important for improving the throughput while maintaining generation accuracy. KVSWAP provides an API to automatically obtain the best group size, but users can configure it to prioritize generation accuracy or throughput.

KV reuse and throughput. To analyse the characteristics and benefits of KV reuse, we evaluate 100 randomly sampled inputs: 50 from the QMSum meeting summarization dataset and 50 from the MuSiQue multi-document QA dataset, covering different use scenarios. Tab. 8 reports reuse rates and throughput across disk types and datasets. Reuse rates remain high across all cases, ranging from 75.3% to 81.2%. The throughput (in tokens/s) achieves

$2.0\times$ – $2.1\times$ speedups on NVMe and $3.8\times$ – $4.0\times$ on eMMC compared to the no-reuse baseline. This indicates that our reuse buffer can effectively improve the throughput, especially for slower disk. Similar trends hold across the evaluated workloads: reuse rates vary little across inputs (standard deviation $\leq 1.1\%$), and throughput variation remains bounded within 2.0 standard deviations on NVMe and 0.6 on eMMC, contributing to no more than 5.6% of the average throughput. These results also justify using the average reuse rate during offline tuning (Appendix A.1).

Number of selected entries. Fig. 13b shows normalized throughput versus the number of selected KV entries MG (omitting the number of KV heads H_{kv}). Increasing MG improves accuracy on NVMe and eMMC but reduces throughput. Beyond $MG = 400$, accuracy gains become marginal while throughput continues to drop, making $MG = 400$ the balanced trade-off and our default setting.

Rolling buffer. We compare generation accuracy with and without the rolling buffer (RB) on KV entry group sizes (G) 2, 4, 8 and 12 (see also Tab. 3 of Appendix B). Disabling RB causes a sharp accuracy drop of $\geq 29\%$, since new KV entries are generated token by token during decoding and often cannot immediately form a group for prediction. The RB allows newly generated entries to participate in computation in a timely manner, helping to preserve accuracy.

Layer-to-Layer approximation. Since the predictor approximates the current-layer input using the previous layer’s input, its accuracy depends on adjacent-layer similarity. Measuring cosine similarity of hidden states on LongBench with Llama3-8B shows it increases with depth: from 25% (Layer 0–1) to 67% (Layer 1–2), then from 75% to 83%, and exceeding 90% in deeper layers. Thus, we keep exact inputs only for the first two layers, preserving accuracy while enabling computation–I/O overlap.

6.2 Overhead Analysis

Latency breakdown. Fig. 13a reports the decoding latency of a single Transformer block in Llama3-8B (batch=8, CL=32K). In FlexGen (FG), loading all KV entries from disk makes I/O the bottleneck. InfiniGen* (IG*) alleviates this with selective loading and our head aggregation scheme, but I/O latency still dominates. KVSWAP without reuse (Ours wo/ reu) better utilizes disk bandwidth, reducing latency by 1.5 $\times$, already comparable to InfiniGen*+reuse (IG*+reu). With reuse enabled, I/O latency of KVSWAP further drops by 4.3 $\times$, incurring only 1.0 ms reuse overhead, and achieves the lowest total latency of 6.9 ms. Apart from the 1.0 ms reuse operation, the main computation (Att+FFN) in KVSWAP takes 2.2 ms, the predictor 1.1 ms, and disk I/O 2.6 ms. Since the predictor and main computation run asynchronously, the reported predictor latency reflects the portion not overlapped with the main computation, likely due to compute resource contention under batched workloads.

Breakdown of memory usage. When applying KVSWAP to Llama3-8B, under a total KV budget of 310 MiB (Tab. 1), the low-rank K cache occupies 256 MiB, corresponding to a compression ratio of $\sigma = 8$. The reuse buffer uses 50 MiB, corresponding to $MG = 400$ selected entries, while the remaining 4 MiB is allocated to the rolling buffer and the CPU staging buffer. Under a tighter budget of 120 MiB, the K cache is further compressed to 64 MiB with $\sigma = 32$, while the other components remain largely unchanged.

6.3 Compare with In-memory Compression

In-memory compression methods [42, 84] compress only the prompt KV while keeping decoding KV fully in memory, making them inefficient for long-context generation, whereas KVSWARE remains memory-efficient. In this scenario, the uncompressed decoding KV memory can easily exceed the compressed prompt KV, largely offsetting the benefits of prompt compression.

To further illustrate this limitation, we compare KVSWARE with the representative in-memory compression method SnapKV [42] using Qwen3-8B on LongBench-Write [16]. We filter out samples with context shorter than 4K, resulting in 59 evaluation samples. This task can produce up to ~30K decoding tokens, creating substantial decoding KV memory. We set the KV memory budget to 350 MiB per batch for KVSWARE, whereas SnapKV requires ~4.2 GiB KV memory per batch due to uncompressed decoding KV.

Under this setting, KVSWARE achieves an average score of 4.28/5, compared to 4.14 for SnapKV, despite using significantly less KV memory. This result highlights the advantage of KVSWARE over in-memory compression methods for long-context generation.

7 Discussion

Naturally, there is room for improvement and further work. We discuss a few points here.

LM weight optimization. KVSWARE optimizes KV cache memory footprint. Weight-centric methods such as quantization [26, 44] and parameter offloading [11, 73] are orthogonal and complementary.

Low-bit KV. KVSWARE is designed to preserve generation quality, and can be combined with low-bit KV compression [45, 48]. For comparison, 2-bit quantization yields an 8× reduction, while KVSWARE achieves more than 30×.

Prefilling optimization. KVSWARE can extend to prefilling by (1) computing attention only for important tokens, and (2) storing only critical KV entries to disk. It also complements prefill-oriented frameworks [34].

OS paging. Using OS-level paging is inefficient, as it is workload-agnostic and misses optimizations like selective KV loading, reduced transfer, and overlapped prefetching.

ShadowKV+reuse. While employing reuse to ShadowKV [60] may improve throughput, it cannot meet the tight on-device memory limits. By algorithm and system co-design, KVSWARE improves both memory efficiency and system throughput.

8 Related Work

KVSWARE builds on several prior foundations, but differs qualitatively from each.

Sparse attention. Sparse attention [20] reduces the memory and compute overhead of attention and KV cache. Static methods apply fixed patterns to select KV entries [17, 21, 36, 80], reducing computation but risking loss of long-range dependencies and accuracy. Dynamic methods adapt token selection to the input, focusing on tokens that dominate attention scores [12, 38, 59, 62, 66, 70, 75, 79, 81, 82, 84]. KVSWARE builds on this insight by identifying important KV entries at each step, but adapts it for on-device disk offloading. Unlike sparse attention, which assumes the KV cache is fully in memory, KVSWARE selectively preloads critical KV groups from disk into memory using two techniques: a compressed K cache

(Sec. 3.2) to reduce memory overhead, and a grouped KV prediction algorithm (Sec. 3.3) to improve disk I/O.

KV-cache offloading and optimization. InfiniGen [41] and ShadowKV [60] move KV entries between the GPU and CPU, while SpeCache applies low-bit compression and speculative prefetching [35]. llm-offload [69] and ArkVale [18] adopt page-based KV managers to recall evicted tokens. These *GPU-CPU offloading* methods assume high-bandwidth, low-latency PCIe connections, whereas mobile and embedded devices rely on NVMe, UFS, eMMC, or SD storage with orders of magnitude lower bandwidth and higher latency. This makes disk-based offloading a fundamentally harder problem. KVSWARE is the first framework designed to tackle this challenge. Optimizations are also proposed for *time-to-first-token (TTFT)* to reduce the startup delay. CacheBlend fuses cached knowledge for faster retrieval-augmented generation [76], IMPRESS [19] accelerates prefix KV loading on server-class SSDs, and CacheGen [46] compresses KV for faster transmission. In contrast, KVSWARE targets iterative decoding on resource-constrained devices, where decoding throughput - not TTFT - is the primary bottleneck.

On-device inference. General-purpose LM serving frameworks such as llama.cpp [2] and MLC-LLM [4] enable inference on consumer hardware. Recent work further improves efficiency through quantization [26, 44, 71], weight optimization [25, 32, 67], contextual sparsity [47], weight offloading [11, 73], model specialization [77], speculative decoding [72], and mobile-oriented model redesigns [49, 50, 63, 64]. Most of these approaches optimize *model weights*. In contrast, KVSWARE targets the growing *KV-cache* overhead by enabling sparse-attention-based KV-cache offloading on mobile devices, addressing a key bottleneck in long-context inference. It complements weight-centric methods to support long-context LM inference on memory-limited devices.

9 Conclusion

We have presented KVSWARE, a new KV-cache offloading scheme for long-context on-device language model inference on mobile and embedded devices. In contrast to prior KV-cache offloading methods designed for server-grade hardware with high memory bandwidth, KVSWARE specifically targets resource-constrained devices and uses non-volatile secondary storage as the swapping medium for KV-cache data. A key contribution of KVSWARE is the co-design of new KV cache management algorithms with a hardware-aware runtime system. This enables substantial efficiency gains over existing KV-cache offloading approaches while preserving generation quality. The framework also demonstrates strong scalability across varying context lengths, batch sizes, model sizes, and tight memory budgets. Through the open-source release of the code and experimental artefact, we hope KVSWARE will provide a foundation for future research and a practical route to deploying powerful long-context language models on everyday devices.

Acknowledgments

This work was supported in part by the UK Engineering and Physical Sciences Research Council (EPSRC) under grant agreement EP/X037304/1 and a Royal Society Industry Fellowship. For any correspondence of this work, please contact Chunwei Xia (Email: C.Xia@leeds.ac.uk) and Zheng Wang (Email: z.wang5@leeds.ac.uk).

References

- [1] NVIDIA. 2024. *Jetson-Orin*. NVIDIA. <https://www.nvidia.com/en-us/autonomous-machines/embedded-systems/jetson-orin/>
- [2] Georgi Gerganov. 2024. *llama.cpp*. Georgi Gerganov. <https://github.com/ggerganov/llama.cpp>
- [3] MediaTek. 2025. *MediaTek Dimensity 9400 Plus*. MediaTek. <https://www.mediatek.com/products/smartphones/mediatek-dimensity-9400-plus>
- [4] MLC Team. 2024. *MLC-LLM*. MLC Team. <https://github.com/mlc-ai/mlc-llm>
- [5] NVIDIA. 2025. *NVIDIA Nsight Systems*. NVIDIA. <https://developer.nvidia.com/nsight-systems>
- [6] NVIDIA. 2025. *NVIDIA Tools Extension Library*. NVIDIA. <https://github.com/NVIDIA/NVTX>
- [7] Orange Pi. 2025. *Orange Pi 5 Pro*. Orange Pi. <http://www.orangepi.org/html/hardWare/computerAndMicrocontrollers/details/Orange-Pi-5-Pro.html>
- [8] Ultralytics. 2025. *Ultralytics YOLO*. Ultralytics. <https://github.com/ultralytics/ultralytics>
- [9] Adobe Inc. 2025. *Adobe Acrobat AI Assistant: Generative AI Tool for PDF Summarization and Smart Q&A*. <https://www.adobe.com/acrobat/generative-ai-pdf.html>
- [10] Joshua Ainslie, James Lee-Thorp, Michiel de Jong, Yury Zemlyanskiy, Federico Lebrón, and Sumit Sanghai. 2023. Gqa: Training generalized multi-query transformer models from multi-head checkpoints. *arXiv preprint arXiv:2305.13245* (2023).
- [11] Keivan Alizadeh, Iman Mirzadeh, Dmitry Belenko, Karen Khatamifard, Minsik Cho, Carlo C Del Mundo, Mohammad Rastegari, and Mehrdad Farajtabar. 2023. Llm in a flash: Efficient large language model inference with limited memory. *arXiv preprint arXiv:2312.11514* (2023).
- [12] Sotiris Anagnostidis, Dario Pavlo, Luca Biggio, Lorenzo Noci, Aurelien Lucchi, and Thomas Hofmann. 2023. Dynamic context pruning for efficient and interpretable autoregressive transformers. In *Proceedings of the 37th International Conference on Neural Information Processing Systems* (New Orleans, LA, USA) (NIPS '23). Curran Associates Inc., Red Hook, NY, USA, Article 2845, 22 pages.
- [13] Apple Inc. 2025. *Apple Intelligence*. <https://apple.com/apple-intelligence>
- [14] Shuai Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Sibao Song, Kai Dang, Peng Wang, Shijie Wang, Jun Tang, et al. 2025. Qwen2. 5-vl technical report. *arXiv preprint arXiv:2502.13923* (2025).
- [15] Yushi Bai, Xin Lv, Jiajie Zhang, Hongchang Lyu, Jiankai Tang, Zhidian Huang, Zhengxiao Du, Xiao Liu, Aohan Zeng, Lei Hou, et al. 2023. Longbench: A bilingual, multitask benchmark for long context understanding. *arXiv preprint arXiv:2308.14508* (2023).
- [16] Yushi Bai, Jiajie Zhang, Xin Lv, Linzhi Zheng, Siqi Zhu, Lei Hou, Yuxiao Dong, Jie Tang, and Juanzi Li. 2024. Longwriter: Unleashing 10,000+ word generation from long context llms. *arXiv preprint arXiv:2408.07055* (2024).
- [17] Iz Beltagy, Matthew E. Peters, and Arman Cohan. 2020. Longformer: The Long-Document Transformer. *CoRR abs/2004.05150* (2020). [arXiv:2004.05150](https://arxiv.org/abs/2004.05150) <https://arxiv.org/abs/2004.05150>
- [18] Renze Chen, Zhuofeng Wang, Beiquan Cao, Tong Wu, Size Zheng, Xiuhong Li, Xuechao Wei, Shengen Yan, Meng Li, and Yun Liang. 2024. ArkVale: Efficient Generative LLM Inference with Recalable Key-Value Eviction. In *Advances in Neural Information Processing Systems 38: Annual Conference on Neural Information Processing Systems 2024, NeurIPS 2024, Vancouver, BC, Canada, December 10 - 15, 2024*, Amir Globersons, Lester Mackey, Danielle Belgrave, Angela Fan, Ulrich Paquet, Jakub M. Tomczak, and Cheng Zhang (Eds.). http://papers.nips.cc/paper_files/paper/2024/hash/cd4b49379efac6e84186a3f3ce108c37-Abstract-Conference.html
- [19] Weijian Chen, Shuibing He, Haoyang Qu, Ruidong Zhang, Siling Yang, Ping Chen, Yi Zheng, Baoxing Huai, and Gang Chen. 2025. {IMPRESS}: An {Importance-Informed}{Multi-Tier} Prefix {KV} Storage System for Large Language Model Inference. In *23rd USENIX Conference on File and Storage Technologies (FAST 25)*. 187–201.
- [20] Rewon Child, Scott Gray, Alec Radford, and Ilya Sutskever. 2019. Generating Long Sequences with Sparse Transformers. *CoRR abs/1904.10509* (2019). [arXiv:1904.10509](https://arxiv.org/abs/1904.10509) <http://arxiv.org/abs/1904.10509>
- [21] Jiayu Ding, Shuming Ma, Li Dong, Xingxing Zhang, Shaohan Huang, Wenhui Wang, Nanning Zheng, and Furu Wei. 2023. LongNet: Scaling Transformers to 1,000,000,000 Tokens. *arXiv:2307.02486 [cs.CL]* <https://arxiv.org/abs/2307.02486>
- [22] Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783* (2024).
- [23] Carl Eckart and Gale Young. 1936. The approximation of one matrix by another of lower rank. *Psychometrika* 1, 3 (1936), 211–218.
- [24] Fireflies.ai Corp. 2025. *Fireflies.ai: AI Meeting Assistant for transcription, summaries, highlights and search*. <https://fireflies.ai/>
- [25] Elias Frantar and Dan Alistarh. 2023. Sparsegpt: Massive language models can be accurately pruned in one-shot. In *International conference on machine learning*. PMLR, 10323–10337.
- [26] Elias Frantar, Saleh Ashkboos, Torsten Hoeftler, and Dan Alistarh. 2022. Gptq: Accurate post-training quantization for generative pre-trained transformers. *arXiv preprint arXiv:2210.17323* (2022).
- [27] Google. 2025. *Google Gemini in Gmail: AI-powered Email Summaries and Natural-Language Q&A*. <https://blog.google/products/gmail/>
- [28] Google. 2025. *Google Lens: Natural-language video and voice search*. <https://lens.google/>
- [29] Google. 2025. *Google Recorder: On-device audio transcription and summarization*. <https://recorder.google.com/>
- [30] Gabriel Haas and Viktor Leis. 2023. What modern nvme storage can do, and how to exploit it: High-performance i/o for high-performance storage engines. *Proceedings of the VLDB Endowment* 16, 9 (2023), 2090–2102.
- [31] Cheng-Ping Hsieh, Simeng Sun, Samuel Kriman, Shantanu Acharya, Dima Rekeshe, Fei Jia, Yang Zhang, and Boris Ginsburg. 2024. RULER: What's the Real Context Size of Your Long-Context Language Models? *arXiv preprint arXiv:2404.06654* (2024).
- [32] Yen-Chang Hsu, Ting Hua, Sungen Chang, Qian Lou, Yilin Shen, and Hongxia Jin. 2022. Language model compression with weighted low-rank factorization. *arXiv preprint arXiv:2207.00112* (2022).
- [33] Naman Jain, King Han, Alex Gu, Wen-Ding Li, Fanjia Yan, Tianjun Zhang, Sida Wang, Armando Solar-Lezama, Koushik Sen, and Ion Stoica. 2024. Livecodebench: Holistic and contamination free evaluation of large language models for code. *arXiv preprint arXiv:2403.07974* (2024).
- [34] Huiqiang Jiang, Yucheng Li, Chengruidong Zhang, Qianhui Wu, Xufang Luo, Surin Ahn, Zhenhua Han, Amir H Abdi, Dongsheng Li, Chin-Yew Lin, et al. 2024. Minference 1.0: Accelerating pre-filling for long-context llms via dynamic sparse attention. *Advances in Neural Information Processing Systems* 37 (2024), 52481–52515.
- [35] Shibo Jie, Yehui Tang, Kai Han, Zhi-Hong Deng, and Jing Han. 2025. SpeCache: Speculative Key-Value Caching for Efficient Generation of LLMs. *arXiv:2503.16163 [cs.CL]* <https://arxiv.org/abs/2503.16163>
- [36] Hongye Jin, Xiaotian Han, Jingfeng Yang, Zhimeng Jiang, Zirui Liu, Chia-Yuan Chang, Huiyuan Chen, and Xia Hu. 2024. LLM Maybe LongLM: SelfExtend LLM context window without tuning. In *Proceedings of the 41st International Conference on Machine Learning* (Vienna, Austria) (ICML '24). JMLR.org, Article 888, 16 pages.
- [37] Greg Kamradt. 2023. *Needle in a Haystack - Pressure Testing LLMs*. https://github.com/gkamradt/LLMTest_NeedleInAHaystack
- [38] Nikita Kitaev, Lukasz Kaiser, and Anselm Levskaya. 2020. Reformer: The Efficient Transformer. In *8th International Conference on Learning Representations, ICLR 2020, Addis Ababa, Ethiopia, April 26–30, 2020*. OpenReview.net. <https://openreview.net/forum?id=rkgNKKHtVb>
- [39] Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph Gonzalez, Hao Zhang, and Ion Stoica. 2023. Efficient memory management for large language model serving with pagedattention. In *Proceedings of the 29th Symposium on Operating Systems Principles*. 611–626.
- [40] Md Tahmid Rahman Laskar, Xue-Yong Fu, Cheng Chen, and Shashi Bhushan Tn. 2023. Building real-world meeting summarization systems using large language models: A practical perspective. *arXiv preprint arXiv:2310.19233* (2023).
- [41] Wonbeom Lee, Jungi Lee, Junghwan Seo, and Jaewoong Sim. 2024. {InfiniGen}: Efficient generative inference of large language models with dynamic {KV} cache management. In *18th USENIX Symposium on Operating Systems Design and Implementation (OSDI 24)*. 155–172.
- [42] Yuhong Li, Yingbing Huang, Bowen Yang, Bharat Venkitesh, Acyr Locatelli, Hanchen Ye, Tianle Cai, Patrick Lewis, and Deming Chen. 2024. Snapkv: Llm knows what you are looking for before generation. *Advances in Neural Information Processing Systems* 37 (2024), 22947–22970.
- [43] Hunter Lightman, Vineet Kosaraju, Yuri Burda, Harrison Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. 2023. Let's verify step by step. In *The twelfth international conference on learning representations*.
- [44] Ji Lin, Jiaming Tang, Haotian Tang, Shang Yang, Wei-Ming Chen, Wei-Chen Wang, Guangxuan Xiao, Xingyu Dang, Chuang Gan, and Song Han. 2024. AWQ: Activation-aware Weight Quantization for On-Device LLM Compression and Acceleration. *Proceedings of Machine Learning and Systems* 6 (2024), 87–100.
- [45] Yujun Lin, Haotian Tang, Shang Yang, Zhekai Zhang, Guangxuan Xiao, Chuang Gan, and Song Han. 2024. Qserve: W4a8kv4 quantization and system co-design for efficient llm serving. *arXiv preprint arXiv:2405.04532* (2024).
- [46] Yuhua Liu, Hanchen Li, Yihua Cheng, Siddhant Ray, Yuyang Huang, Qizheng Zhang, Kuntai Du, Jiayi Yao, Shan Lu, Ganesh Ananthanarayanan, et al. 2024. Cachegen: Kv cache compression and streaming for fast large language model serving. In *Proceedings of the ACM SIGCOMM 2024 Conference*. 38–56.
- [47] Zichang Liu, Jue Wang, Tri Dao, Tianyi Zhou, Binhang Yuan, Zhao Song, Anshumali Shrivastava, Ce Zhang, Yundong Tian, Christopher Re, et al. 2023. Deja vu: Contextual sparsity for efficient llms at inference time. In *International Conference on Machine Learning*. PMLR, 22137–22176.
- [48] Zirui Liu, Jiayi Yuan, Hongye Jin, Shaochen Zhong, Zhaozhuo Xu, Vladimir Braverman, Beidi Chen, and Xia Hu. 2024. Kivi: A tuning-free asymmetric 2bit

- quantization for kv cache. *arXiv preprint arXiv:2402.02750* (2024).
- [49] Zechun Liu, Changsheng Zhao, Forrest Iandola, Chen Lai, Yuandong Tian, Igor Fedorov, Yunyang Xiong, Ernie Chang, Yangyang Shi, Raghuraman Krishnamoorthi, et al. 2024. Mobilellm: Optimizing sub-billion parameter language models for on-device use cases. *arXiv preprint arXiv:2402.14905* (2024).
 - [50] Sachin Mehta, Mohammad Hossein Sekhavat, Qingqing Cao, Maxwell Horton, Yanzi Jin, Chenfan Sun, Seyed Iman Mirzadeh, Mahyar Najibi, Dmitry Belenko, Peter Zatloukal, et al. 2024. Openelm: An efficient language model family with open training and inference framework. In *Workshop on Efficient Systems for Foundation Models II@ ICML2024*.
 - [51] Stephen Merity, Caiming Xiong, James Bradbury, and Richard Socher. 2016. Pointer sentinel mixture models. *arXiv preprint arXiv:1609.07843* (2016).
 - [52] Jayashree Mohan, Rohan Kadekodi, and Vijay Chidambaram. 2017. Analyzing IO amplification in Linux file systems. *arXiv preprint arXiv:1707.08514* (2017).
 - [53] Otter.ai, Inc. 2025. Otter.ai: AI Meeting Agent for real-time transcription, summarization, and action-item extraction. <https://otter.ai/>.
 - [54] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, et al. 2019. Pytorch: An imperative style, high-performance deep learning library. *Advances in neural information processing systems* 32 (2019).
 - [55] Reiner Pope, Sholto Douglas, Aakanksha Chowdhery, Jacob Devlin, James Bradbury, Jonathan Heek, Kefan Xiao, Shivani Agrawal, and Jeff Dean. 2023. Efficiently scaling transformer inference. *Proceedings of machine learning and systems* 5 (2023), 606–624.
 - [56] Qualcomm Technologies, Inc. 2024. Snapdragon 8 Elite Mobile Platform. <https://www.qualcomm.com/products/mobile/snapdragon/smartphones/snapdragon-8-series-mobile-platforms/snapdragon-8-elite-mobile-platform>.
 - [57] Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J Liu. 2020. Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of machine learning research* 21, 140 (2020), 1–67.
 - [58] Ying Sheng, Lianmin Zheng, Binhang Yuan, Zhuohan Li, Max Ryabinin, Beidi Chen, Percy Liang, Christopher Ré, Ion Stoica, and Ce Zhang. 2023. Flexgen: High-throughput generative inference of large language models with a single gpu. In *International Conference on Machine Learning*. PMLR, 31094–31116.
 - [59] Prajwal Singhan, Siddharth Singh, Shwai He, Soheil Feizi, and Abhinav Bhatt. 2024. Loki: Low-rank keys for efficient sparse attention. *arXiv preprint arXiv:2406.02542* (2024).
 - [60] Hanshi Sun, Li-Wen Chang, Wenlei Bao, Size Zheng, Ningxin Zheng, Xin Liu, Harry Dong, Yuejie Chi, and Beidi Chen. 2024. Shadowkv: Kv cache in shadows for high-throughput long-context llm inference. *arXiv preprint arXiv:2410.21465* (2024).
 - [61] Jiaming Tang, Yilong Zhao, Kan Zhu, Guangxuan Xiao, Baris Kasicki, and Song Han. 2024. Quest: Query-aware sparsity for efficient long-context llm inference. *arXiv preprint arXiv:2406.10774* (2024).
 - [62] Jiaming Tang, Yilong Zhao, Kan Zhu, Guangxuan Xiao, Baris Kasicki, and Song Han. 2024. QUEST: query-aware sparsity for efficient long-context LLM inference. In *Proceedings of the 41st International Conference on Machine Learning (Vienna, Austria) (ICML'24)*. JMLR.org, Article 1955, 11 pages.
 - [63] Gemma Team, Morgane Riviere, Shreya Pathak, Pier Giuseppe Sessa, Cassidy Hardin, Surya Bhupatiraju, Léonard Hussenot, Thomas Mesnard, Bobak Shahriari, Alexandre Ramé, et al. 2024. Gemma 2: Improving open language models at a practical size. *arXiv preprint arXiv:2408.00118* (2024).
 - [64] Omkar Thawakar, Ashmal Vayani, Salman Khan, Hisham Cholakkal, Rao M Anwer, Michael Felsberg, Tim Baldwin, Eric P Xing, and Fahad Shahbaz Khan. 2024. Mobillama: Towards accurate and lightweight fully transparent gpt. *arXiv preprint arXiv:2402.16840* (2024).
 - [65] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need. *Advances in neural information processing systems* 30 (2017).
 - [66] Hanrui Wang, Zhekai Zhang, and Song Han. 2021. SpAtten: Efficient Sparse Attention Architecture with Cascade Token and Head Pruning. In *2021 IEEE International Symposium on High-Performance Computer Architecture (HPCA)*. 97–110. doi:10.1109/HPCA51647.2021.00018
 - [67] Xin Wang, Yu Zheng, Zhongwei Wan, and Mi Zhang. 2024. Svd-llm: Truncation-aware singular value decomposition for large language model compression. *arXiv preprint arXiv:2403.07378* (2024).
 - [68] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. 2022. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems* 35 (2022), 24824–24837.
 - [69] Jianbo Wu, Jie Ren, Shuangyan Yang, Konstantinos Parasyris, Giorgos Georgakoudis, Ignacio Laguna, and Dong Li. 2025. LM-Offload: Performance Model-Guided Generative Inference of Large Language Models with Parallelism Control. In *2025 IEEE International Parallel and Distributed Processing Symposium Workshops (IPDPSW)*. 840–849. doi:10.1109/IPDPSW66978.2025.00134
 - [70] Wei Wu, Zhuoshi Pan, Chao Wang, Liyi Chen, Yunchu Bai, Tianfu Wang, Kun Fu, Zheng Wang, and Hui Xiong. 2025. TokenSelect: Efficient Long-Context Inference and Length Extrapolation for LLMs via Dynamic Token-Level KV Cache Selection. *arXiv:2411.02886 [cs.CL]* <https://arxiv.org/abs/2411.02886>
 - [71] Guangxuan Xiao, Ji Lin, Mickael Seznec, Hao Wu, Julien Demouth, and Song Han. 2023. Smoothquant: Accurate and efficient post-training quantization for large language models. In *International Conference on Machine Learning*. PMLR, 38087–38099.
 - [72] Daliang Xu, Wangsong Yin, Xin Jin, Ying Zhang, Shiyun Wei, Mengwei Xu, and Xuanzhe Liu. 2023. Llmcad: Fast and scalable on-device large language model inference. *arXiv preprint arXiv:2309.04255* (2023).
 - [73] Zhenliang Xue, Yixin Song, Zeyu Mi, Le Chen, Yubin Xia, and Haibo Chen. 2024. PowerInfer-2: Fast Large Language Model Inference on a Smartphone. *arXiv preprint arXiv:2406.06282* (2024).
 - [74] An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, et al. 2025. Qwen3 technical report. *arXiv preprint arXiv:2505.09388* (2025).
 - [75] Songlin Wang, Bailin Wang, Yikang Shen, Rameswar Panda, and Yoon Kim. 2024. Gated linear attention transformers with hardware-efficient training. In *Proceedings of the 41st International Conference on Machine Learning (Vienna, Austria) (ICML'24)*. JMLR.org, Article 2333, 23 pages.
 - [76] Jiayi Yao, Hanchen Li, Yuhuan Liu, Siddhant Ray, Yihua Cheng, Qizheng Zhang, Kuntai Du, Shan Lu, and Junchen Jiang. 2025. CacheBlend: Fast large language model serving for RAG with cached knowledge fusion. In *Proceedings of the Twentieth European Conference on Computer Systems*. 94–109.
 - [77] Rongjie Yi, Liwei Guo, Shiyun Wei, Ao Zhou, Shuangwang Wang, and Mengwei Xu. 2023. Edgemoe: Fast on-device inference of moe-based large language models. *arXiv preprint arXiv:2308.14352* (2023).
 - [78] Shukang Yin, Chaoyou Fu, Sirui Zhao, Ke Li, Xing Sun, Tong Xu, and Enhong Chen. 2024. A survey on multimodal large language models. *National Science Review* 11, 12 (2024), nwae403.
 - [79] Jingyang Yuan, Huazuo Gao, Damai Dai, Junyu Luo, Liang Zhao, Zhengyan Zhang, Zhenda Xie, Yuxing Wei, Lean Wang, Zhiping Xiao, Yuqing Wang, Chong Ruan, Ming Zhang, Wenfeng Liang, and Wangding Zeng. 2025. Native Sparse Attention: Hardware-Aligned and Natively Trainable Sparse Attention. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Wanxiang Che, Joyce Nabende, Ekaterina Shutova, and Mohammad Taher Pilehvar (Eds.). Association for Computational Linguistics, Vienna, Austria, 23078–23097. doi:10.18653/v1/2025.acl-long.1126
 - [80] Manzil Zaheer, Guru Guruganesh, Avinava Dubey, Joshua Ainslie, Chris Alberti, Santiago Ontanon, Philip Pham, Anirudh Ravula, Qifan Wang, Li Yang, and Amr Ahmed. 2020. Big bird: transformers for longer sequences. In *Proceedings of the 34th International Conference on Neural Information Processing Systems (Vancouver, BC, Canada) (NIPS '20)*. Curran Associates Inc., Red Hook, NY, USA, Article 1450, 15 pages.
 - [81] Biao Zhang, Ivan Titov, and Rico Sennrich. 2021. Sparse Attention with Linear Units. *arXiv:2104.07012 [cs.CL]*
 - [82] Haojie Zhang, Zhenyu Wu, Zhenyu Zhang, Shixing Liu, Zujie Ren, Jingji Chen, Mingjie Zhan, Zirui Liu, Chen-Yu Lee, Yuchen Zhang, Haichuan Yang, Yuxiang Liu, Yafan He, and Zhaofeng He. 2024. SemSA: Semantic Sparse Attention is hidden in Large Language Models.. In *The Twelfth International Conference on Learning Representations*. <https://openreview.net/forum?id=eG9AkHtYYH>
 - [83] Haopeng Zhang, Philip S Yu, and Jiawei Zhang. 2025. A systematic survey of text summarization: From statistical methods to large language models. *Comput. Surveys* 57, 11 (2025), 1–41.
 - [84] Zhenyu Zhang, Ying Sheng, Tianyi Zhou, Tianlong Chen, Lianmin Zheng, Ruisi Cai, Zhao Song, Yuandong Tian, Christopher Ré, Clark Barrett, et al. 2023. H2o: Heavy-hitter oracle for efficient generative inference of large language models. *Advances in Neural Information Processing Systems* 36 (2023), 34661–34710.
 - [85] Junjie Zhou, Yan Shu, Bo Zhao, Boya Wu, Zhengyang Liang, Shitao Xiao, Minghao Qin, Xi Yang, Yongping Xiong, Bo Zhang, et al. 2025. Mlvu: Benchmarking multi-task long video understanding. In *Proceedings of the Computer Vision and Pattern Recognition Conference*. 13691–13701.
 - [86] Jinguo Zhu, Wei Yun Wang, Zhe Chen, Zhaoyang Liu, Shenglong Ye, Lixin Gu, Hao Tian, Yuchen Duan, Weijie Su, Jie Shao, et al. 2025. InternVL3: Exploring advanced training and test-time recipes for open-source multimodal models. *arXiv preprint arXiv:2504.10479* (2025).