

KASH: 1-BIT KEY-VALUE CACHE QUANTIZATION VIA ASYMMETRIC HASHING

Yifan Zhang, Qinghao Hu, Zhihui Wei, Jian Cheng

Nanjing University of Science and Technology · Institute of Automation, Chinese Academy of Sciences



Presenter: Yifan Zhang



ICPR 2026 · Lyon, France



南京理工大学
NANJING UNIVERSITY OF SCIENCE & TECHNOLOGY



计算机科学与工程学院
School of Computer Science and Engineering



中国科学院自动化研究所
Institute of Automation, Chinese Academy of Sciences





Outline

1. Background

2. Our Method

3. Experimental Results

4. Conclusion



Outline

1. Background

2. Our Method

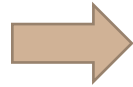
3. Experimental Results

4. Conclusion

1. Background

KV Cache Bottleneck

- Stores a key and value for every token
- Scales with batch size, context length and layers
- **Long-context decoding becomes memory-bound**



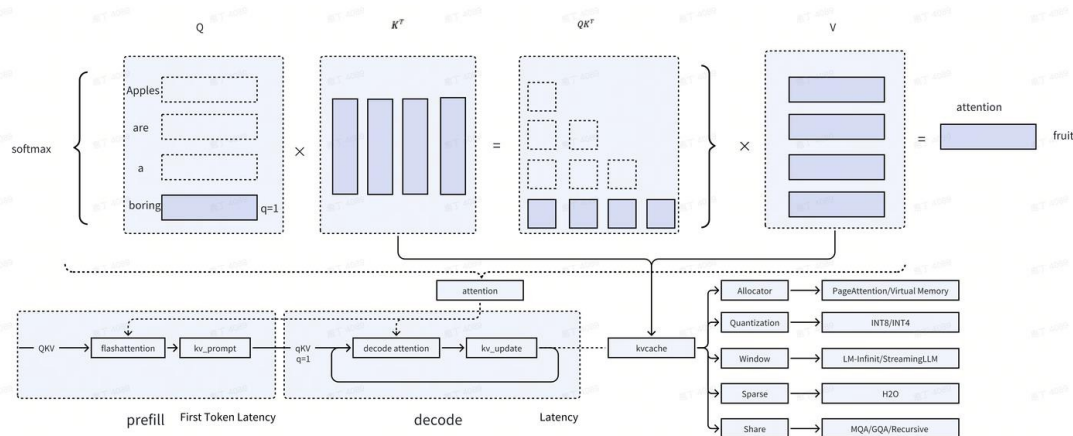
Quantization

- Less cache storage and memory bandwidth
- Grouped N-bit stores extra scales and zero-points
- KIVI-2 therefore averages about 3 bits per element



1-bit Barrier

- Keys are consumed through $q^T k$ retrieval scores
- Sign-only keys perturb $q^T k$ logits
- Softmax amplifies small score errors



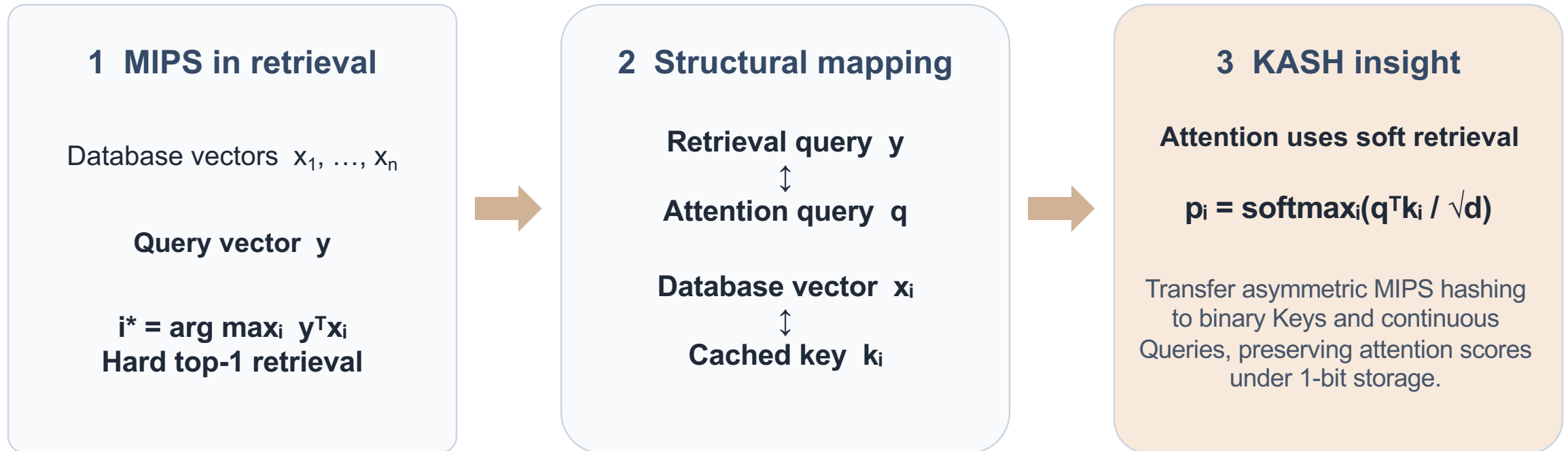
hidden
cost

quality
collapse

We need to push average storage below the ≈ 3 -bit floor while preserving the attention scores $q^T k$.

From MIPS Retrieval to Attention Compression

MIPS(Maximum Inner Product Search) originates in database and information retrieval; attention shares the same inner-product retrieval structure.



KASH transfers asymmetric MIPS hashing from information retrieval to 1-bit KV-cache compression

Why 1-bit Key quantization becomes an inner-product problem?

Attention consumes each key through $q^T k$; coordinate-wise reconstruction is only an indirect proxy.

Conventional objective: reconstruct K

$$\min ||K - \hat{K}||_F^2$$

At 1 bit, $\hat{K}$ keeps mainly the sign pattern and loses magnitude.

A small reconstruction error does not guarantee $Q^T \hat{K} \approx Q^T K$ for future queries.

Optimize what
attention uses



MIPS-inspired objective: preserve scores

$$\min ||B^T Z - K^T Q||_F^2$$

Binary Keys: $B = \text{sign}(H_k^T K) \in \{-1, +1\}$

Continuous Queries: $Z = H_q^T Q \in \mathbb{R}$

Learn B and Z so that binary-key inner products directly approximate the original attention logits.

KASH directly optimizes attention scores under 1-bit Key storage; H_k and H_q are learned separately.



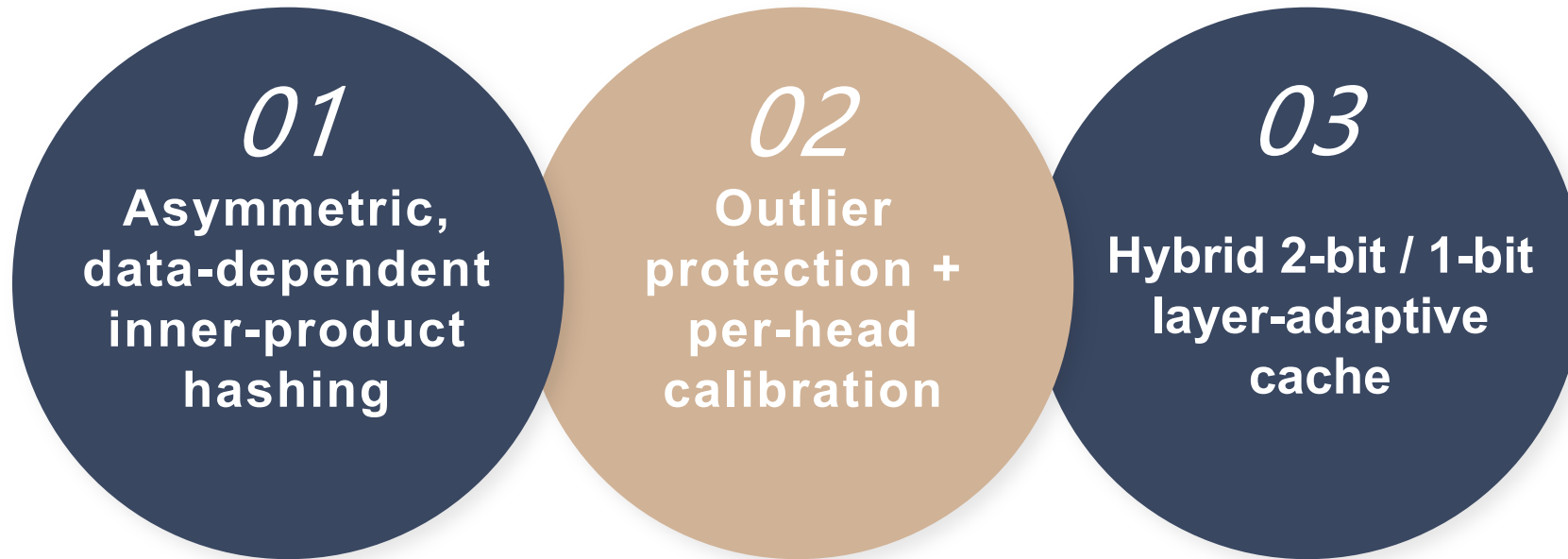
Outline

1. Background

2. Our Method

3. Experimental Results

4. Conclusion



2. Method

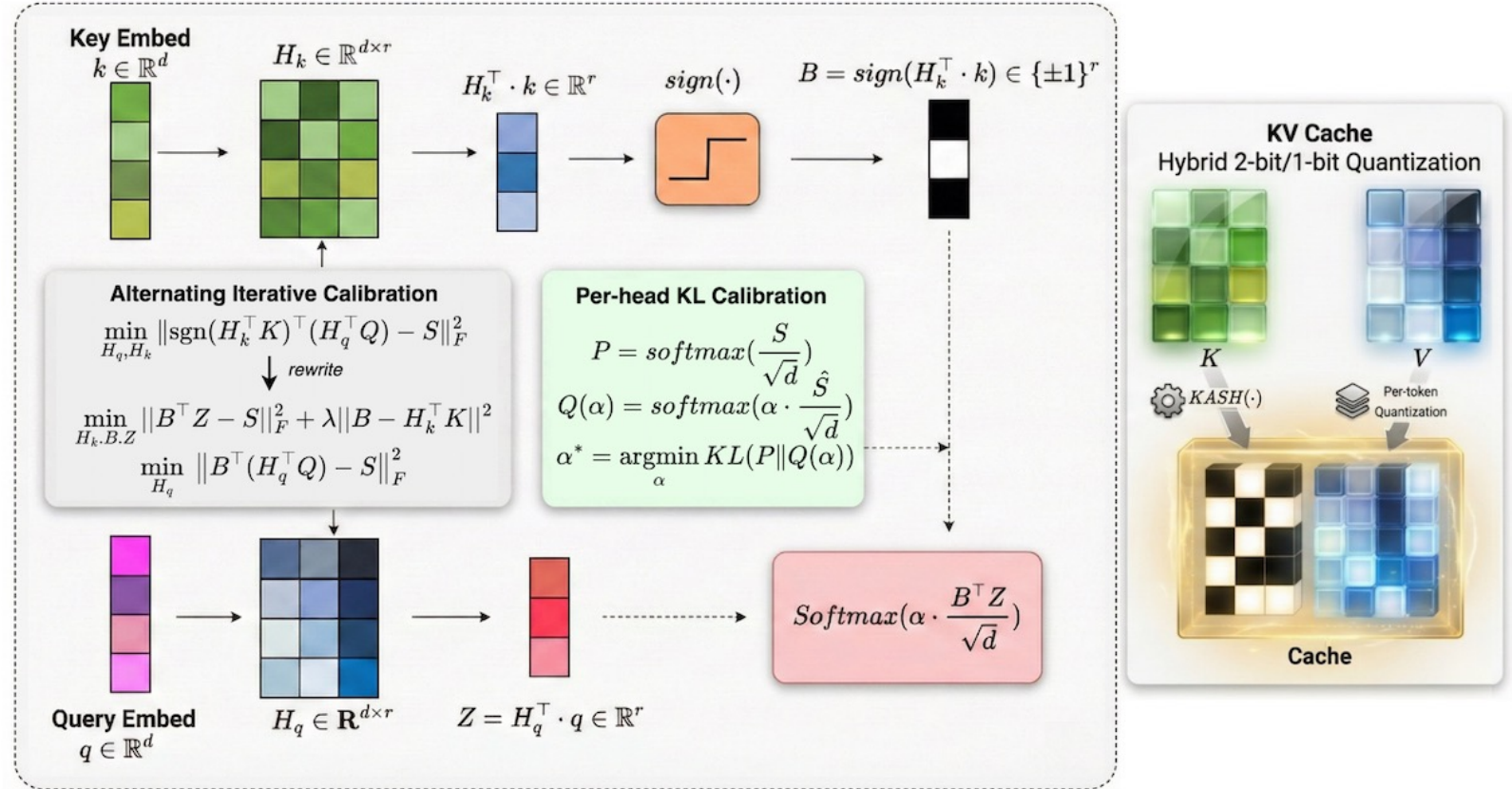
■ Proposed KASH pipeline

➤ Offline calibration:

- Learn H_k and H_q per layer-head
- Detect outlier channels

➤ Online decoding:

- Store bit-packed binary keys
- Keep queries continuous
- Reuse fixed maps at every step



■ Problem Formulation

Asymmetric inner-product fitting keeps Keys discrete and Queries continuous.

1 Define reference scores

Collect calibration Keys K and Queries Q for each layer-head, then form S .

2 Build asymmetric embeddings

Binary Key code B uses $\text{sign}(\cdot)$; projected Query Z remains real-valued.

3 Fit the inner products

Optimize H_k and H_q so $B^T Z$ approximates the original score matrix S .

① Reference inner-product matrix

$$S \triangleq K^\top Q \quad (\text{or } S \triangleq \frac{1}{\sqrt{d}} K^\top Q).$$



② Asymmetric Key-Query mapping

$$B \triangleq \text{sgn}(H_k^\top K) \in \{-1, +1\}^{r \times M}, \quad Z \triangleq H_q^\top Q \in \mathbb{R}^{r \times N}.$$



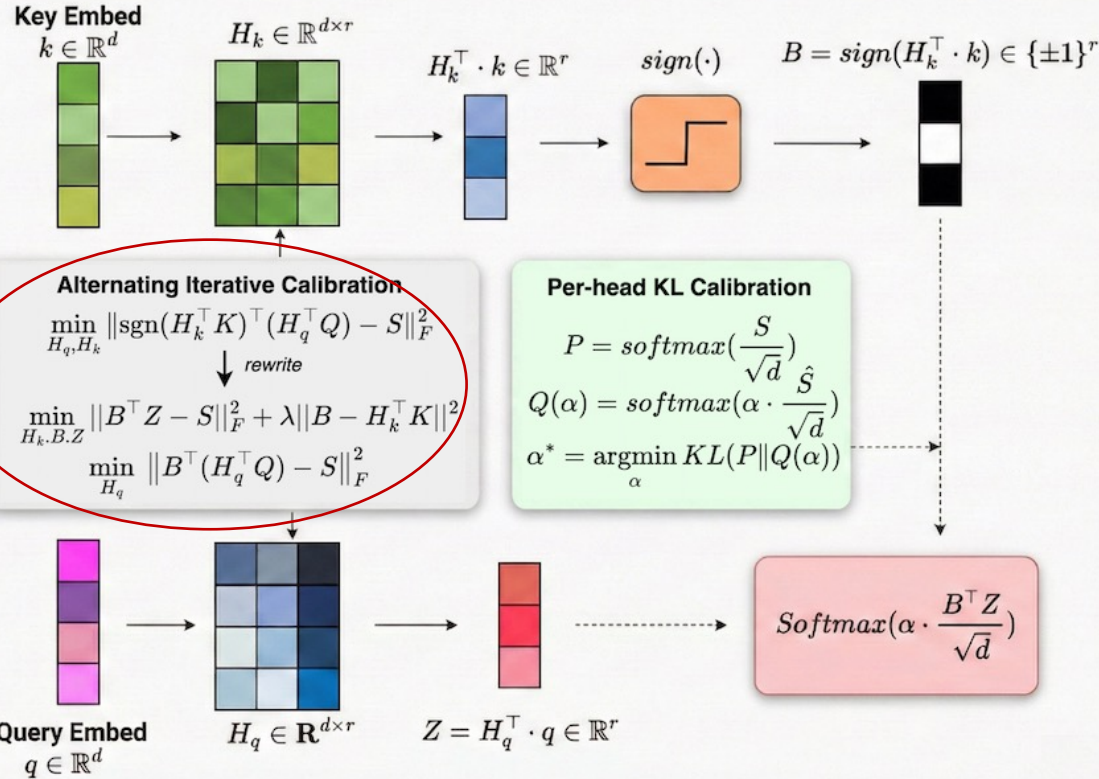
③ Inner-product fitting objective

$$\min_{H_k, H_q} \left\| \text{sgn}(H_k^\top K)^\top (H_q^\top Q) - S \right\|_F^2 \equiv \min_{B, Z} \|B^\top Z - S\|_F^2.$$

1-bit discrete Keys + real-valued projected Queries $\rightarrow$ preserve $q^\top k$

2. Method

■ Alternating optimization



① Fix $H_q \rightarrow Z = H_q^T Q$ is fixed

$$\min_{H_k} \|\text{sgn}(H_k^T K)^T Z - S\|_F^2 \rightarrow \min_{H_k, B} \|B^T Z - S\|_F^2 + \lambda \|B - H_k^T K\|_F^2,$$

② Alternate H_k and B

B fixed $\rightarrow$ update H_k $H_k^* = \arg \min_{H_k} \|B - H_k^T K\|_F^2 \Rightarrow H_k^* = (KK^T)^{-1}KB^T,$

H_k fixed $\rightarrow$ update B $b_t = \text{sgn}(d_t - z_t \tilde{Z}^T \tilde{B}),$

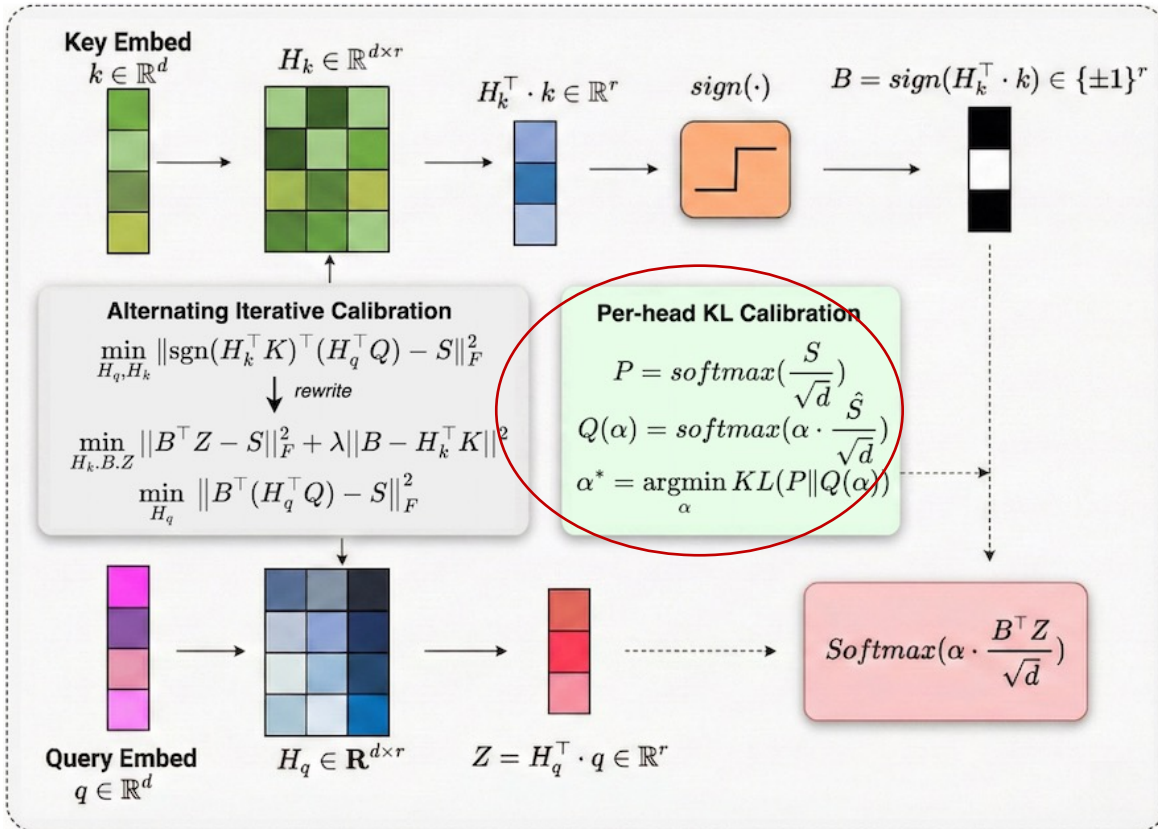
③ Fix $B \rightarrow$ update H_q in closed form

$$\min_{H_q} \|B^T (H_q^T Q) - S\|_F^2 \rightarrow H_q^* = (QQ^T)^{-1}Q S^T B^T (BB^T)^{-1},$$

④ Repeat until converged

2. Method

■ Per-head KL calibration



learn one α for each layer-head

- ① **Reference** $P = \text{softmax}(S / \sqrt{d})$ Full attention
- ② **Calibrated** $Q(\alpha) = \text{softmax}(\alpha \cdot \hat{S} / \sqrt{d})$ Scale the approximate logits
- ③ **Optimize** $\alpha^* = \arg \min \text{KL}(P \| Q(\alpha))$ One-dimensional offline search

Algorithm 1. KASH Key Cache Quantization Framework

PHASE I: CALIBRATION (OFFLINE)

Require: Calibration data $\{(Q, K)\}$, bitwidth r , outlier ratio ρ .

Ensure: Transforms $\{H_q, H_k\}$, scalars $\{\alpha\}$, masks $\{\mathcal{O}\}$.

```

1: for layer  $l \in \mathcal{L}_{\text{late}}$  and head  $h$  do
2:   Detect outliers  $\mathcal{O}$  (top- $\rho$ ); Split  $K \rightarrow [K_{\text{in}}, K_{\text{out}}]$ ; Init  $H_q, H_k$  via PCA.
3:   repeat
4:     Update  $B \leftarrow \text{sgn}(H_k^\top K_{\text{in}})$  and  $H_k$  to align directions.
5:     Update  $H_q \leftarrow \arg \min \|B^\top (H_q^\top Q) - S\|_F^2$  and calibrate  $\alpha$ .
6:   until convergence
7: end for

```

PHASE II: HYBRID DECODING (ONLINE)

```

8: for each decoding step  $n$  and layer  $l$  do
9:   if  $l \notin \mathcal{L}_{\text{late}}$  then  $\triangleright$  Early layers: 2-bit
10:    KIVI.Write( $k_n, v_n$ )
11:   else  $\triangleright$  Late layers: 1-bit asymmetric
12:    Split  $k_n \rightarrow [k_{\text{in}}, k_{\text{out}}]$  via  $\mathcal{O}$ .
13:    Quantize  $b_n \leftarrow \text{sgn}(H_k^\top k_{\text{in}})$ ; Store bit-packed  $b_n$  and FP16  $k_{\text{out}}$ .
14:    Attention Score:  $\hat{s} \leftarrow \text{softmax} \left( \frac{1}{\sqrt{d}} \left( \underbrace{\alpha \cdot (B_{\text{cache}}^\top H_q^\top q_n)}_{\text{1-bit Kernel}} + \underbrace{K_{\text{out}}^\top q_{\text{out}}}_{\text{FP16}} \right) \right)$ .
15:   end if
16: end for

```

Offline calibration

- Detect outliers
- Initialize the two transforms
- Alternate updates of B , H_k , and H_q
- Calibrate alpha

Online decoding

- Early layers: 2-bit KIVI
- Late layers: 1-bit KASH
- Values: 2-bit per-token



Outline

1. Background

2. Our Method

3. Experimental Results

4. Conclusion

3. Experimental Results



■ Setup

➤ Models & Calibration

- LLaMA-2 7B and 13B
- WikiText-2: 128×2048 tokens

➤ Benchmarks

- LongBench: 7 long-context tasks
- LM-Eval: TruthfulQA, CoQA, GSM8K

■ LongBench Results on LLaMA-2-7B

➤ RepoB-P: RepoBench-P · MNews: Multi-News

Method	Avg.bit	TriviaQA	RepoB-P	Qasper	TREC	QMSum	LCC	MNews	Avg.
FP16	16	87.72	59.82	9.52	66	21.28	66.66	3.51	44.93
KIVI-2	3	87.42	60.23	9.31	66	20.5	66.88	1.14	44.50
KVQuant-3	3.33	87.55	59.5	13.67	66	16.53	61.52	1.65	43.77
QJL	3.08	87.47	58.52	9.56	66	20.83	64.27	1.3	43.99
KIVI-2 / KIVI-1	2.5	2.67	16.62	2.59	18	1.49	20.8	0.47	8.95
KASH (Ours)	2.22	87.51	57.76	9.61	66	20.84	64.94	2.5	44.17

3. Experimental Results



■ LongBench and LM-Eval Results on LLaMA-2-13B

➤ RepoB-P: RepoBench-P · MNews: Multi-News

Method	Avg.bit	TriviaQA	RepoB-P	Qasper	TREC	QMSum	LCC	MNews	Avg.
FP16	16	87.87	56.42	9.32	70	21.38	66.61	3.71	45.04
KIVI-2	3	87.31	53.66	8.27	69.5	20.69	65.08	3.19	43.96
QJL	3.08	87.38	53.16	9.18	69.5	20.05	63.34	2.84	43.64
KIVI-2/KIVI-1	2.5	0	13.46	0	19	0	18.61	0.01	7.30
KASH (Ours)	2.22	87.47	52.43	9.83	70	20.69	63.05	2.84	43.76

➤ Standard-context evaluation across TruthfulQA, CoQA, and GSM8K

Method	Avg.Bit	TruthfulQA	CoQA	GSM8K	Avg.
FP16	16	29.53	66.37	22.67	39.52
KIVI-2	3	29.84	66.23	20.77	38.95
KIVI-2/KIVI-1	2.5	29.43	65.57	16.3	37.1
KASH (Ours)	2.2	31.43	66.57	20.22	39.41

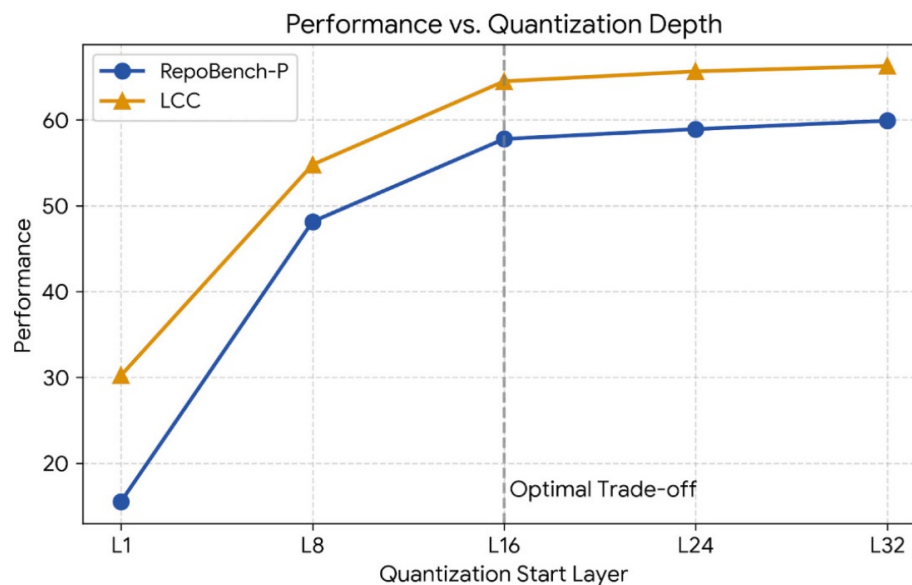
3. Experimental Results

■ Ablation Results

➤ Hq and α

Method	TriviaQA	RepoBench-P	TREC	Avg.
FP16 (no quant.)	87.72	59.82	66	71.18
Ours (Full): Hk $\neq$ Hq, w/ α	87.51	57.76	66	70.42
w/o Hq (set Hq = I), w/ α	84.06	50.59	64	66.22
w/o α (fix $\alpha = 1$), w/ Hq	87.31	56.69	66	70
w/o Hq and w/o α	83.86	47.59	63	64.82

➤ Quantization Depth



4. Conclusion



ultra-long contexts / distribution shifts / value-cache learning / combinations with sparse attention...

THANKS!



Reporter : Yifan Zhang

yifan_zhang@njust.edu.cn



南京理工大学
NANJING UNIVERSITY OF SCIENCE & TECHNOLOGY



中国科学院自动化研究所
Institute of Automation, Chinese Academy of Sciences

