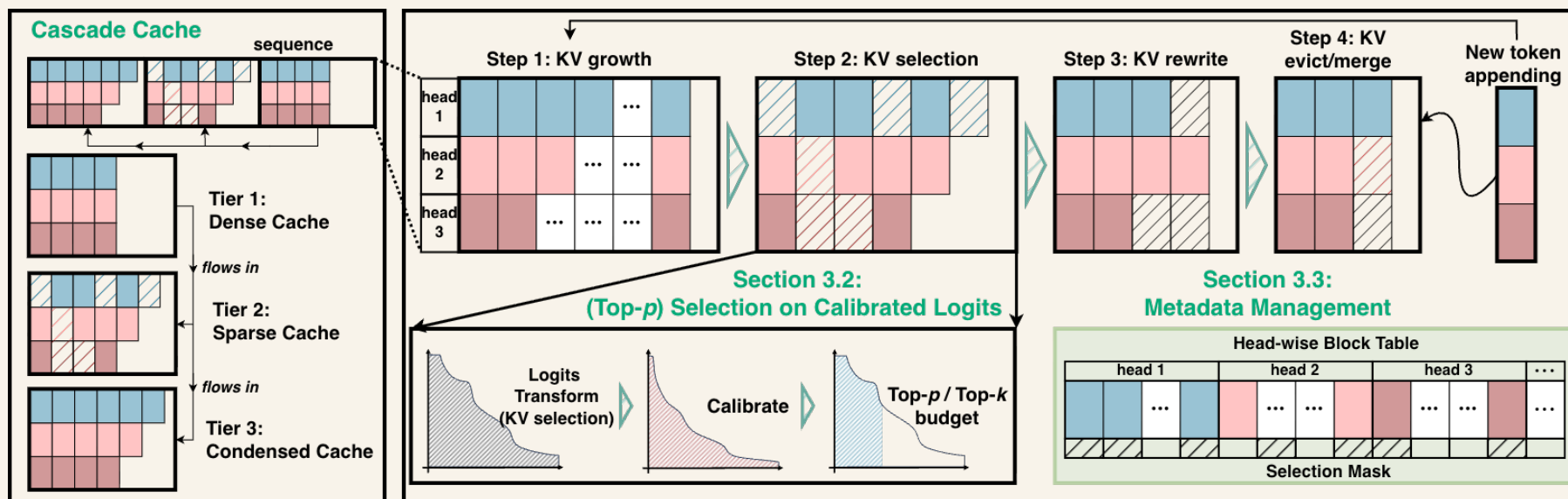




HARD-KV: Head-Adaptive Regularization for Decoding-time KV Compression

Yuxuan Yang¹, Feiyang Ren¹, Bowen Zeng¹, Dalin Zhang², JinPeng Chen³, Gang Chen¹, Huan Li¹



¹The State Key Laboratory of Blockchain and Data Security, Zhejiang University

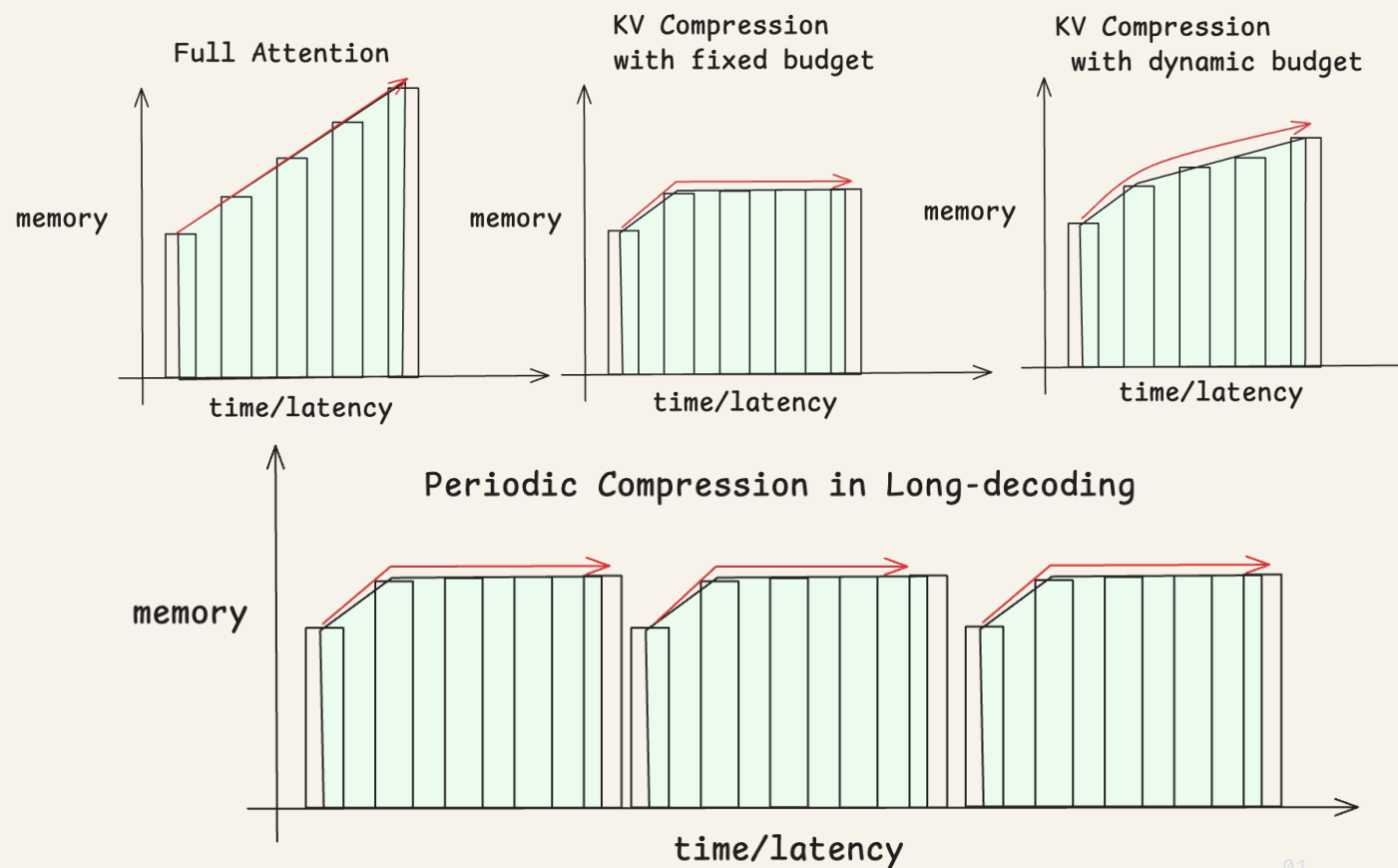
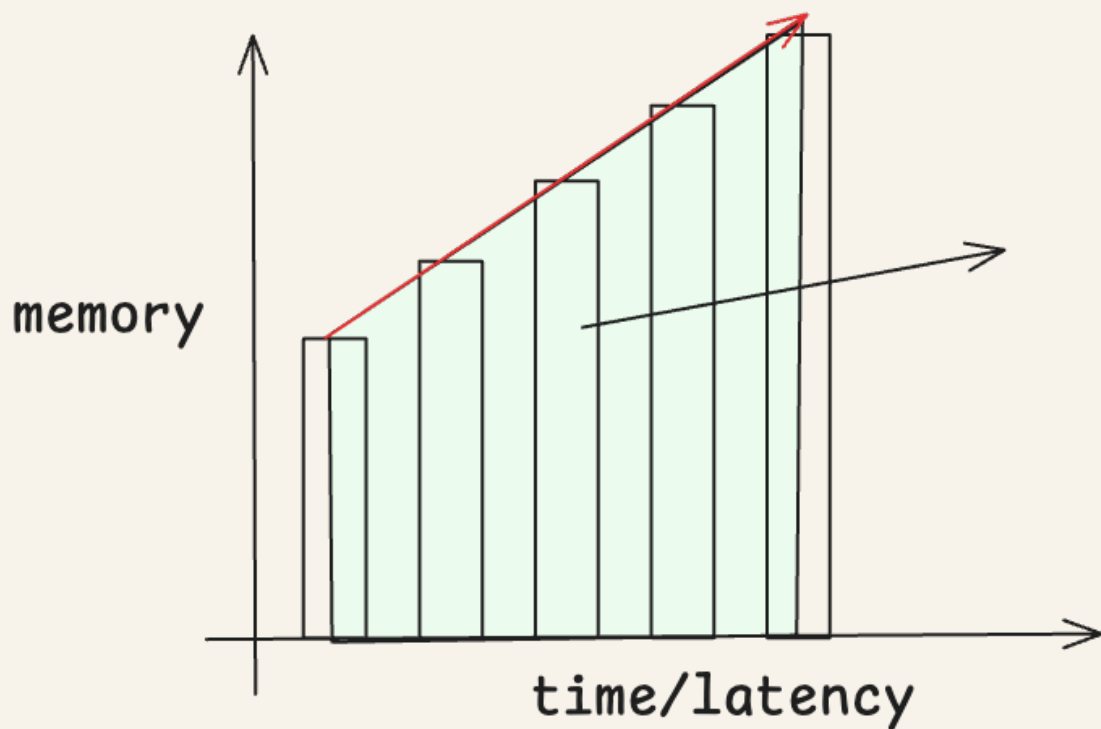
²Space Information Research Institute, Hangzhou Dianzi University

³School of Computer Science (National Pilot Software Engineering School), BUPT

Background: Decoding-time Compression

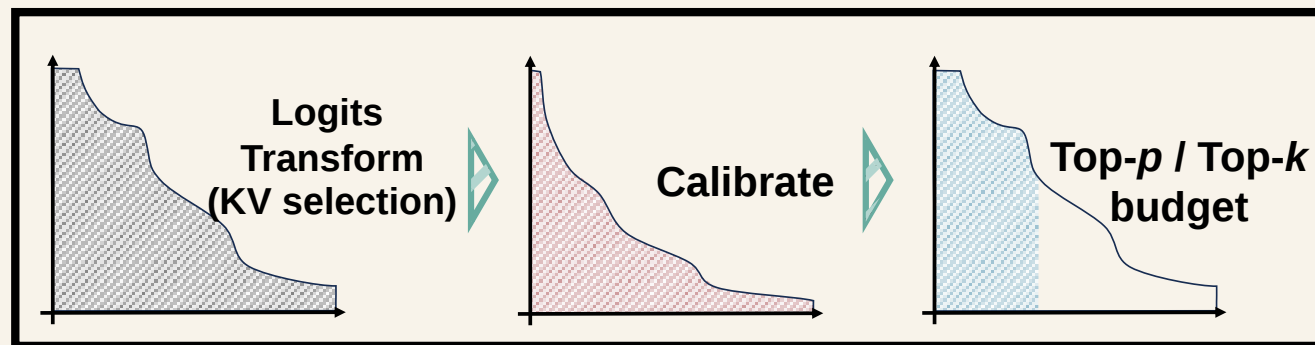
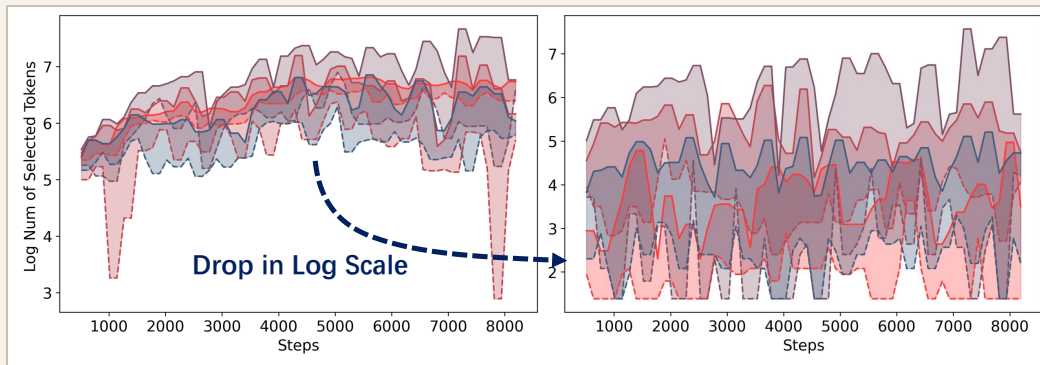
Memory-Time-Integral: the standing-point to think of latency-throughput trade-offs

Memory-Time-Integral (MTI)

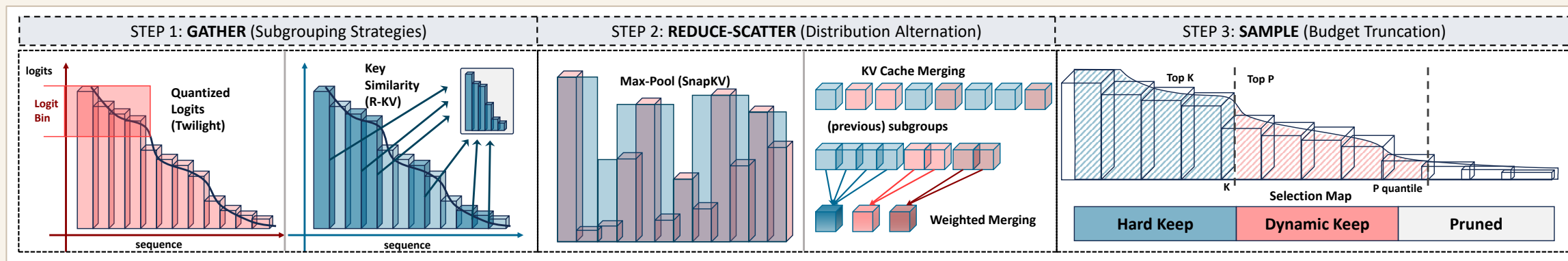


Natural Dynamic Budget: Top-P budget

Selection collapse



If a P mass corresponds to K selections in raw attention, we keep (roughly) same number of selections for different heuristics



Challenge: Static-Dynamic Mismatch

Algorithm side

Head-adaptive Top-p selection lets each head keep the tokens it needs.

Budgets fluctuate by context, layer, and head; this is exactly what protects long reasoning trajectories.

Top-p budget

probability
mass budget

Head-Adaptive

fine-grained
allocation

**STATIC-
DYNAMIC
MISMATCH**

System side

PagedAttention, continuous batching, and CUDA Graphs want stable block layouts.

Irregular per-head indices create metadata work, fragmented loading, and graph-breaking dynamic shapes.

CUDA Graphs

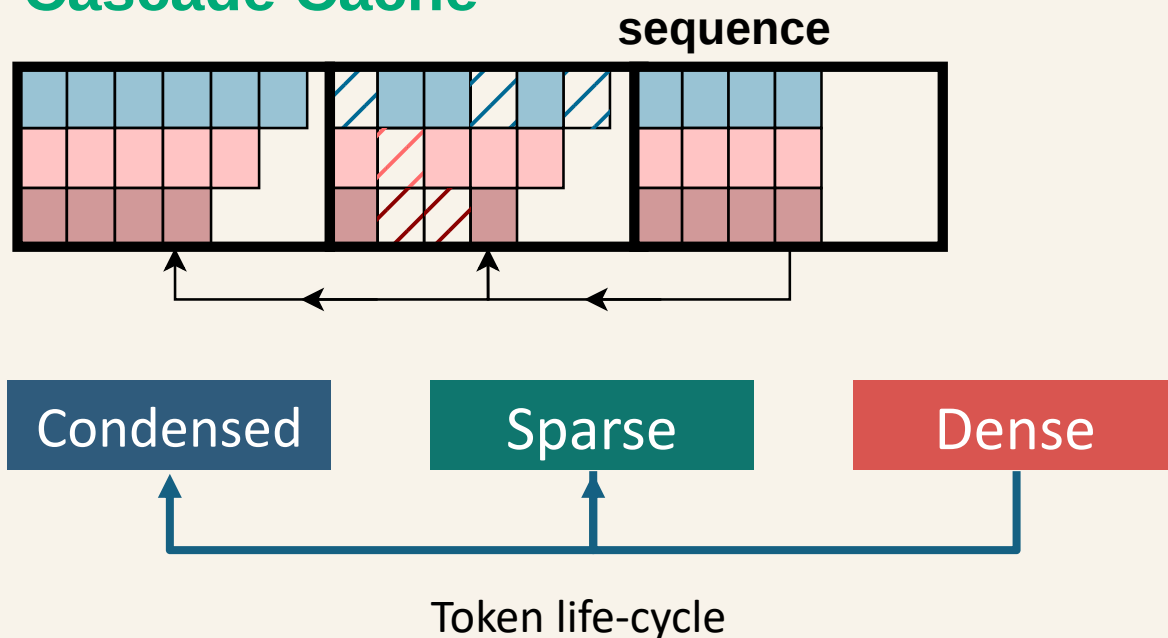
static
execution
shape

Regular Page Index

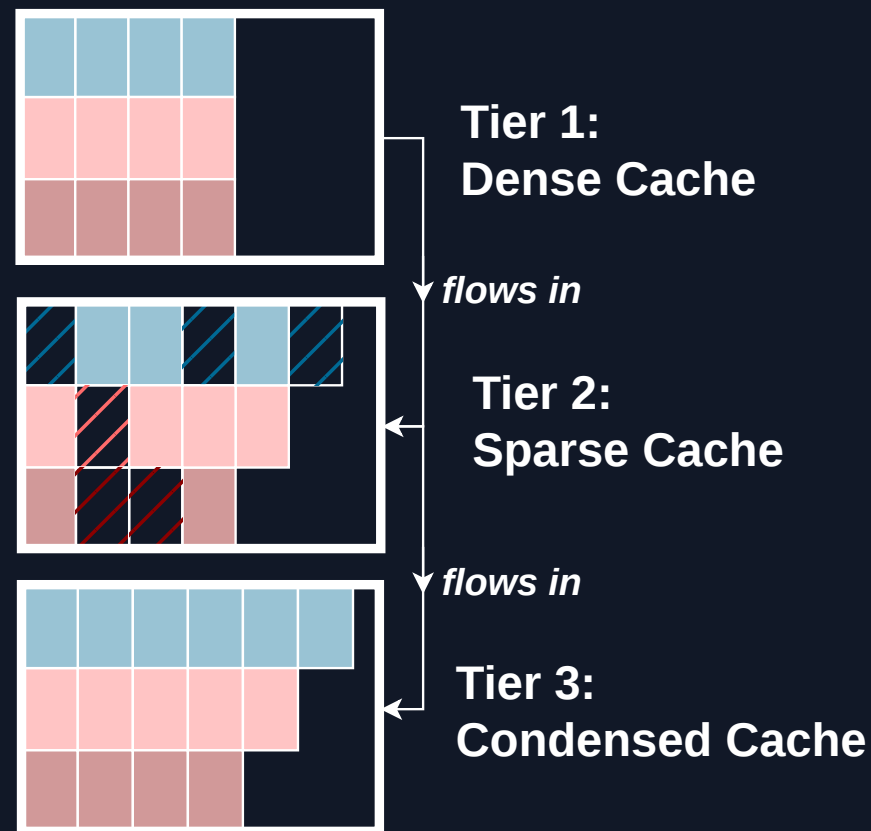
block table
contract

Intuition: Sequential-hybrid of linear and full KV cache

Cascade Cache

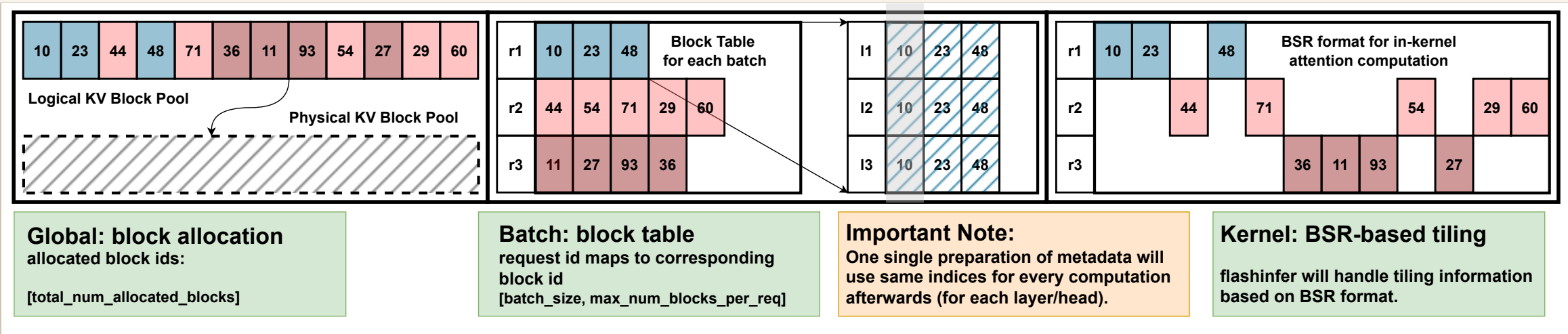


Cascade Cache: Three-tier Contract

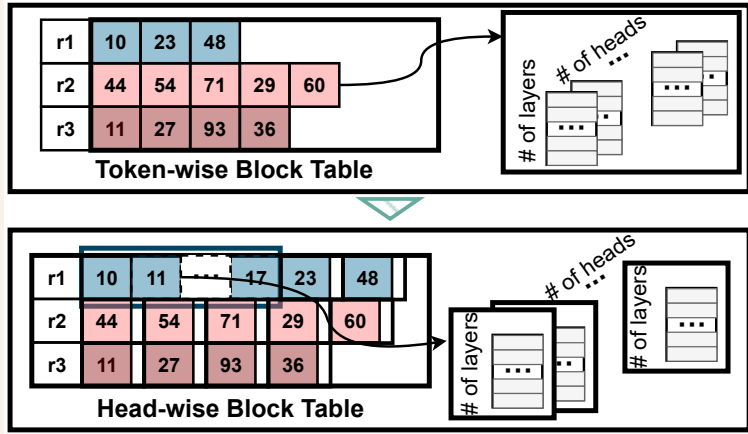


Inference Engine Adaptations

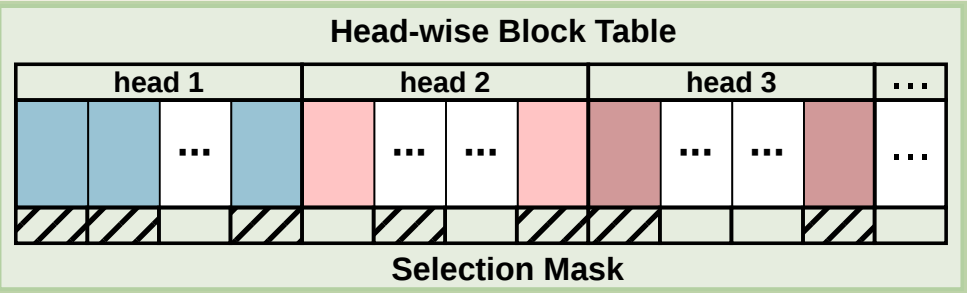
Three metadata levels



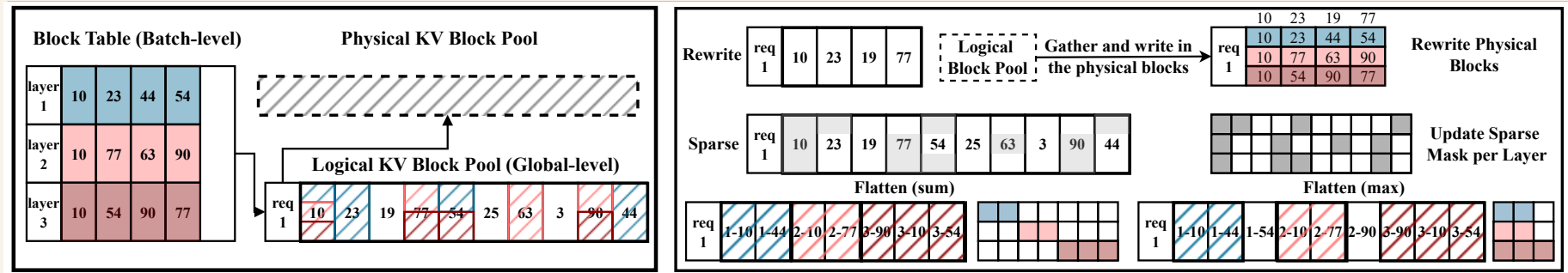
Marginal adjustments with head-wise block table



Supported by FlashInfer kernel library



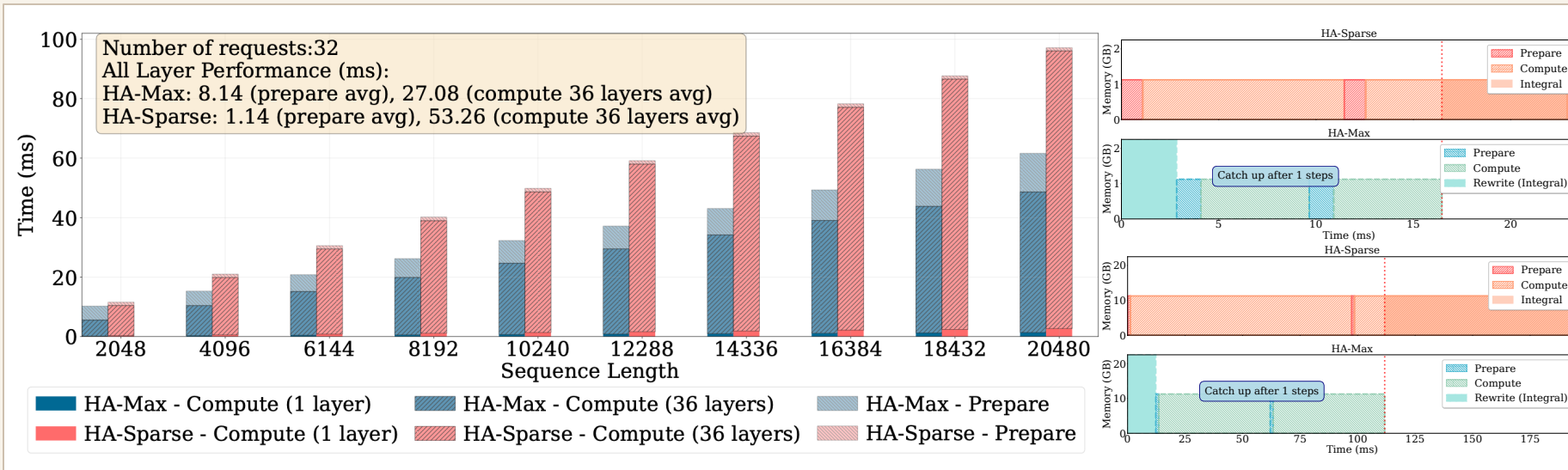
Index Regularization: unifying indices across layers



Rewrite
contiguous cache
read selected blocks, write regularized physical layout

Sparse load
mask at kernel
regularize indices by carrying unused cache positions

Flatten
head/layer index
choose max or sum depending on cost trade-off layout



Rewrite: a naïve but efficient approach
Can catch-up counterparts in 1 step by MTI

Experiment Takeaway

1.5x
performance gain

2x
throughput improvement

90%+
sparsity utilization

Throughput breakdowns under dynamic budget

Base system



+ CUDA Graph



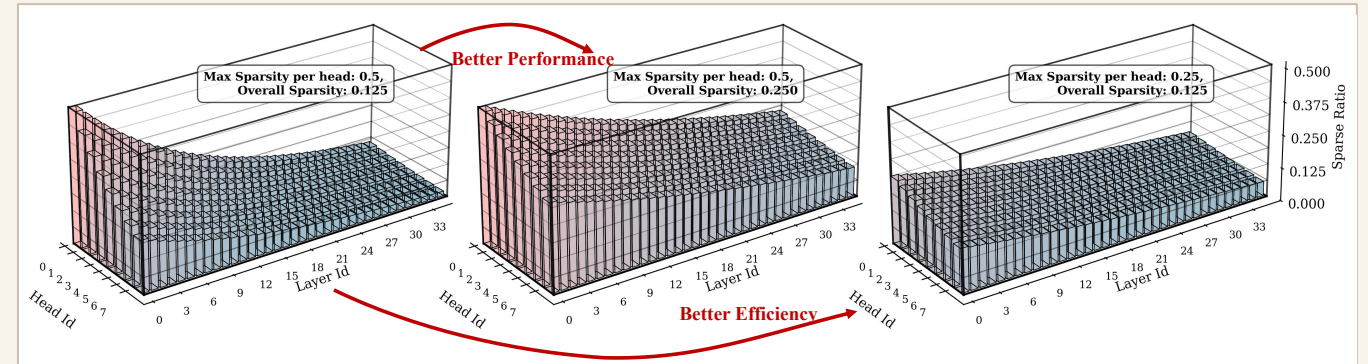
+ Top-k pruning



Naïve top p



+ Ours avg



Takeaways on tuning performances

Use k to cap physical cache growth;
use p to preserve probability mass. Hard-KV makes this two-knob trade-off executable.

Thank you