

Plasma: A Layout-Aware Benchmark Reveals Memory Layout Matters for Graph-based ANNS on GPU

Yutaro Oguri¹ Mai Nishimura² Yusuke Matsui¹

¹The University of Tokyo ²OMRON SINIC X Corporation

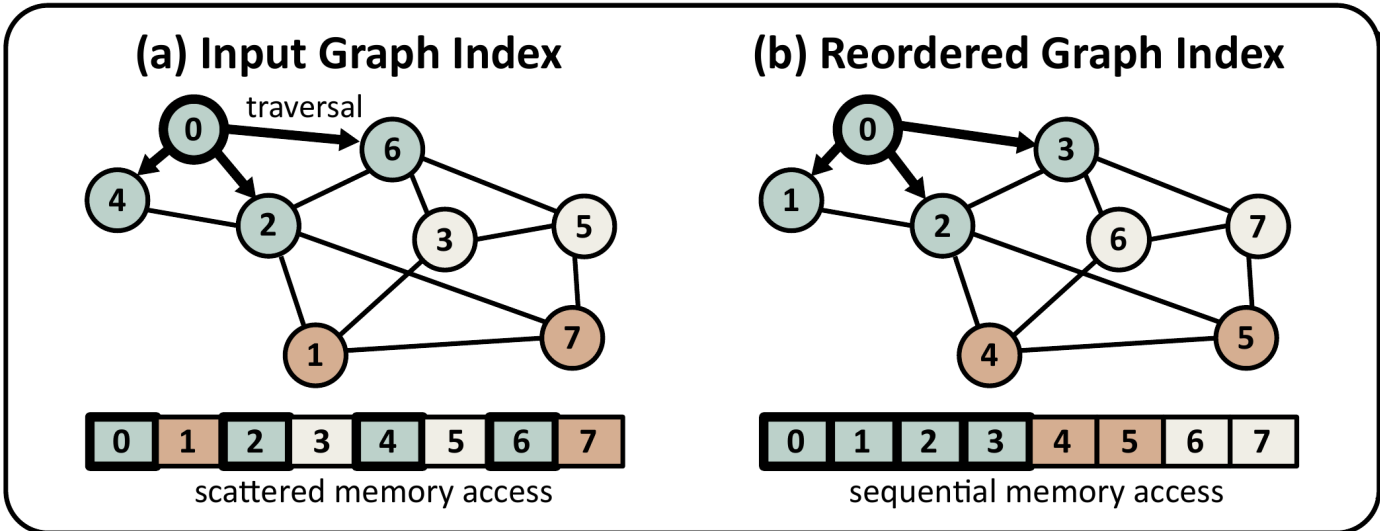


Project Page

1 Graph-based ANNS on GPU

- Core primitive of modern vector search app
- GPU is promising, like CAGRA [Ootomo+, ICDE 2024]
- Bandwidth-bound far more than CPU
- Our scope: **Graph-based x GPU**

2 Memory layout matters



GPU-optimized ANNS speed with enhanced memory utilization

- Effective on CPU execution [Coleman+, NeurIPS 2022]
- Reordering gives co-accessed vertices near IDs
- Scattered access → coalesced
- Proven on CPU, but **still unexplored on GPU**

3 Issue: No fair evaluation platform

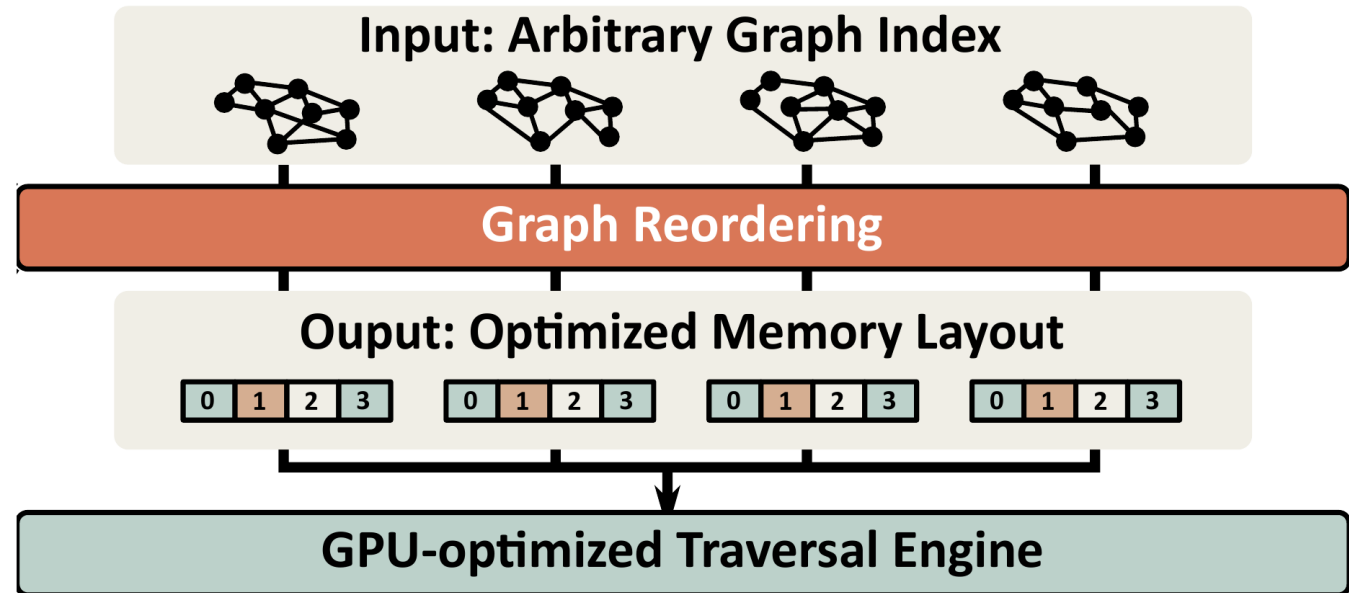
- Each index ships its own search impl. (kernel)
- Topology and implementation are **entangled**
- Existing: focus on, not aware of memory layout

	ANN Benchmark	BigANN	VDB Bench	VIBE	PLASMA
Platform	CPU	CPU	CPU	CPU+GPU	GPU
Modern embedding	—	partial	partial	✓	✓
HW profiling	—	—	—	—	✓
Layout-aware	—	—	—	—	✓
Reproduction	docker	docker	docker	apptainer	pixi

4 Contributions

- PLASMA**: Topology decoupled from layout under one unified GPU kernel for the fair evaluation
- First systematic study**: 4 indices x 12 datasets x 5 orderings
- Surprise**: Graph Topology < Memory Layout reordering alone up to **+80% QPS**
- Attribution**: HW profiling reveals DRAM bandwidth, not L1/L2 hit rate, explains well
- Open release**: pixi build for easy reproduction

5 Our proposal: PLASMA



Fair cross-index comparison without implementation bias

- a Graph adapter**: Input adjacency list format
Enable NSG and Vamana on GPU for the first time
- b Unified kernel**: cuVS-style search kernel
For fair evaluation independent from impl. hack
- c Layout arrangement**: Reorder vertex ids

6 Experiments

Indices	CAGRA, NSG, Vamana, NN-Descent
Datasets	Modern 12 datasets, dim=96-1536
Reordering	Degree Sort, Hub Sort, GOrder, RCM
Hardware	NVIDIA A100 (80GB), A6000 (48GB)
Metrics	Recall@10, QPS

HEADLINE RESULTS

Up to 80%

QPS gain at equal recall

10–30%

typical gain per index–dataset pair

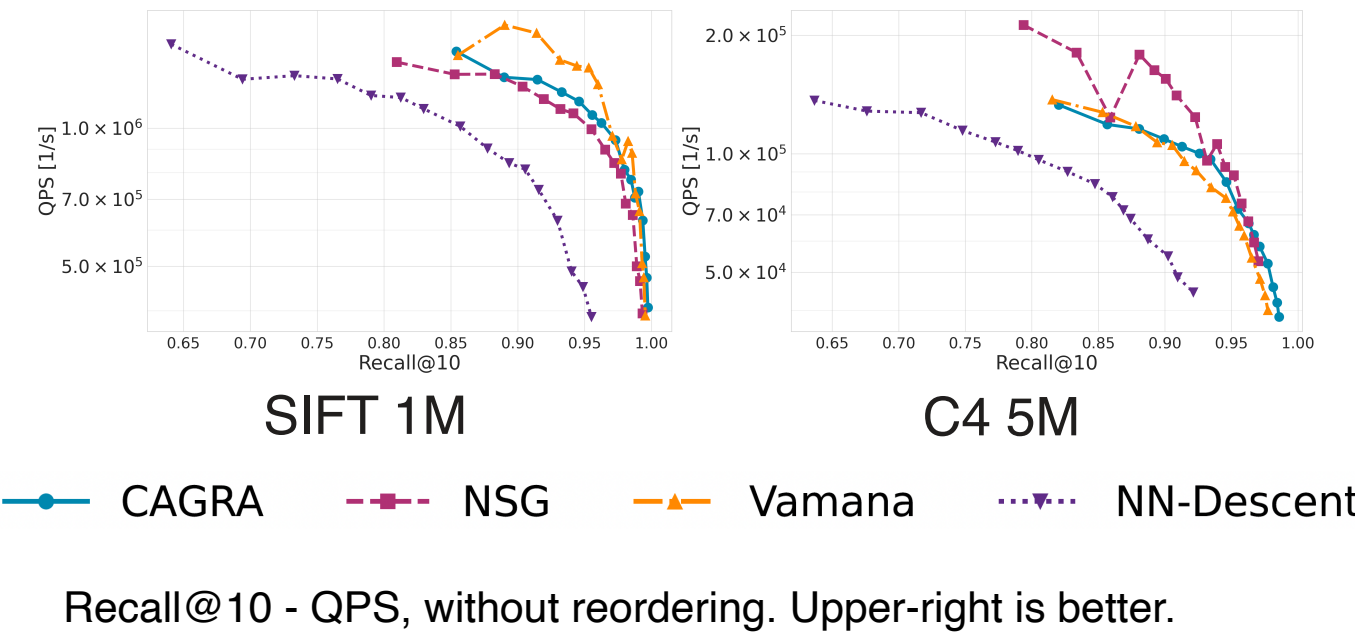
1.2x

CPU-oriented NSG can outperform GPU-native CAGRA

7 Findings

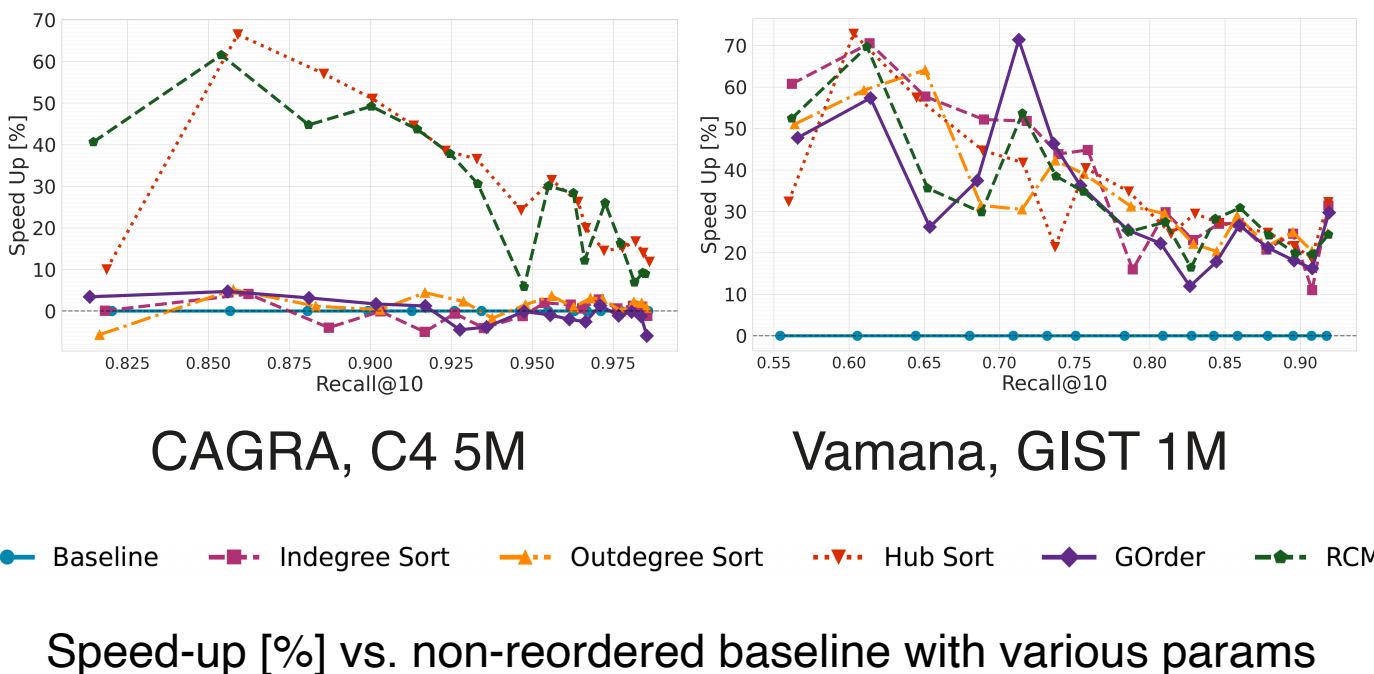
(1) Topology is less significant than expected

- NSG and Vamana can beat CAGRA
*CPU-oriented indices running on GPU!
- NSG accelerated by 1.2x on C4 5M
- Gaps widen on high-dimensional embeddings
- Each indices have largely different topologies



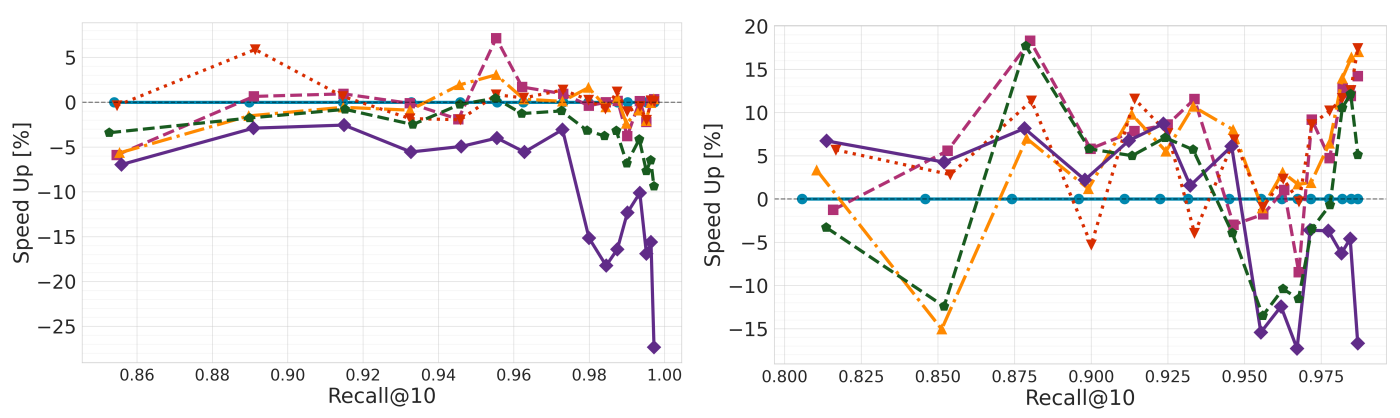
(2) Reordering improves QPS substantially

- QPS improvement: Up to +80%
- No one single winner of reordering methods
- Effectiveness depends on search parameters



Some failure cases

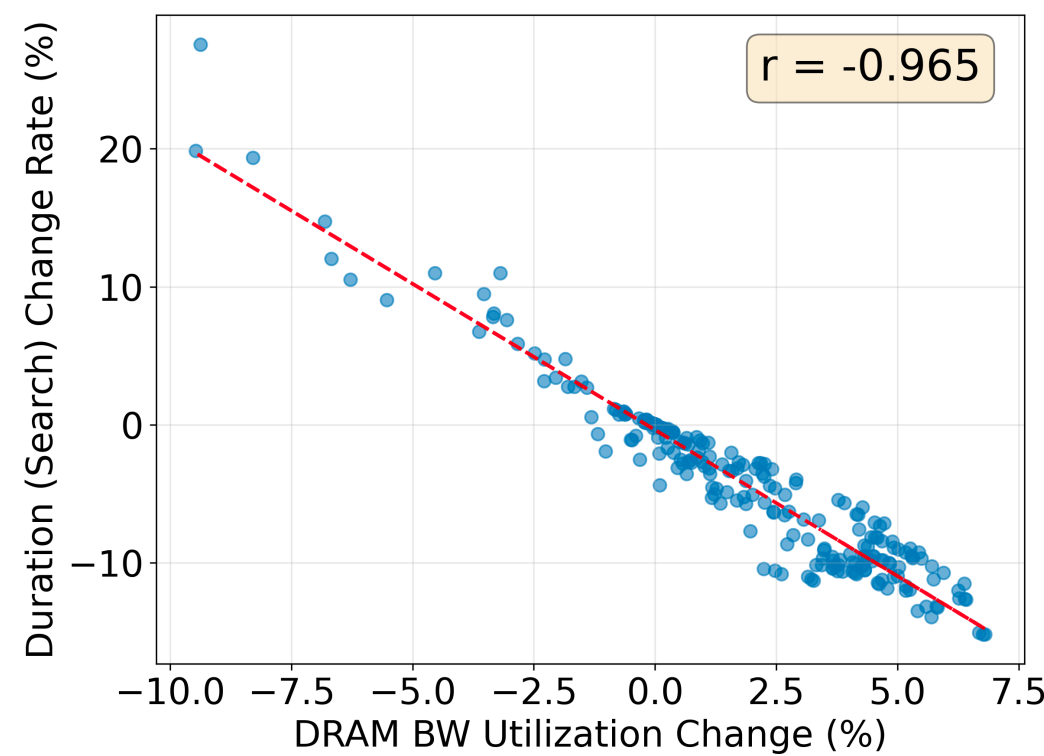
- QPS improvement is unstable
- No one single winner of reordering methods



CAGRA, SIFT 1M NSG, Deep 1M

(3) DRAM bandwidth explains the speed-up

- L1 hit $\approx 0\%$ (0.007 - 0.025%), L2 hit at 8 - 28%
- Cache hit does not explain well
- DRAM bandwidth tracks the speed-up effectively



One point per (index, dataset, reordering) at a certain parameter

- GPU with higher memory bandwidth benefits more

TAKEAWAY

On GPU ANNS, the memory layout of graph matters more than how nodes are connected

- Extend to latest H100 and B100 architectures
- Can lead to memory layout-aware graph index
- Can extend to online reordering scenario