

On Similarity of Computational Kernels in our Codes and Proxies

Workshop on Reproducible and Sustainable HPC Benchmarking | ICS 2026

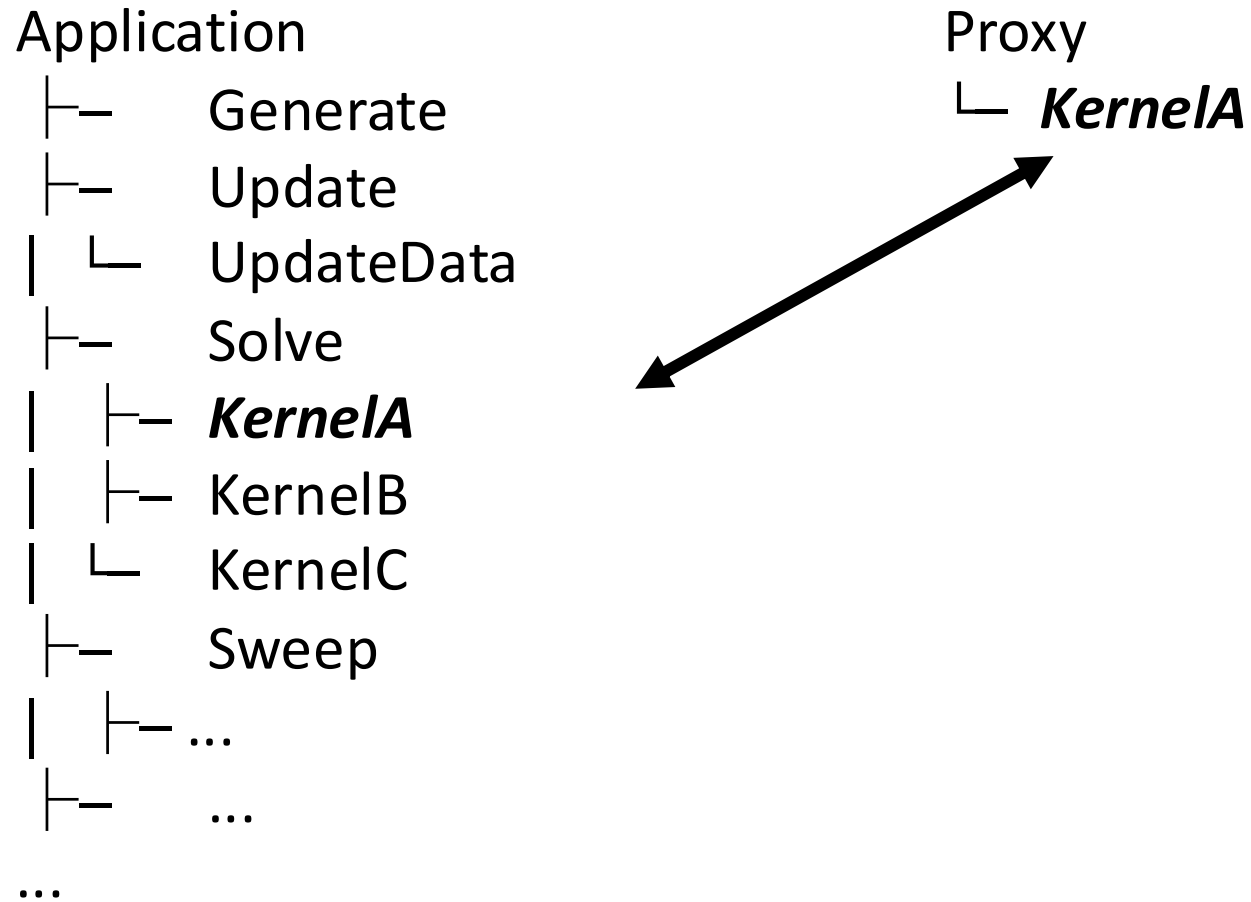


July 6, 2026

Michael McKinsey, Stephanie Brink, **Olga Pearce**
pearce8@llnl.gov



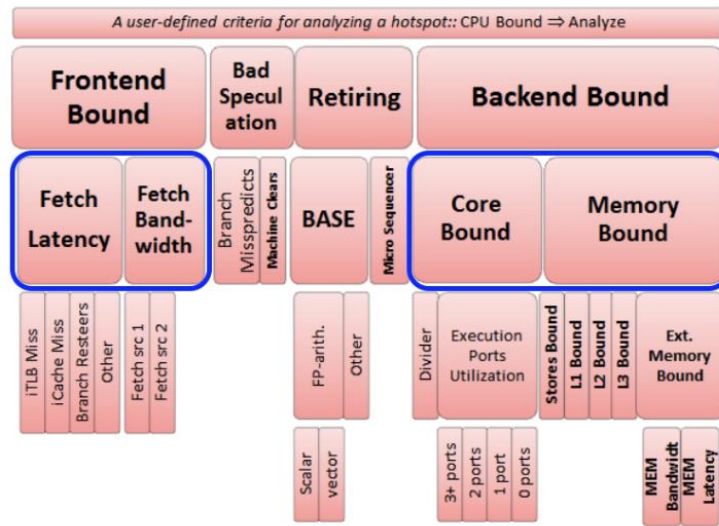
Use benchmarks to evaluate hardware and procure systems



Proxies enable running subsets of computational kernels from complex applications.

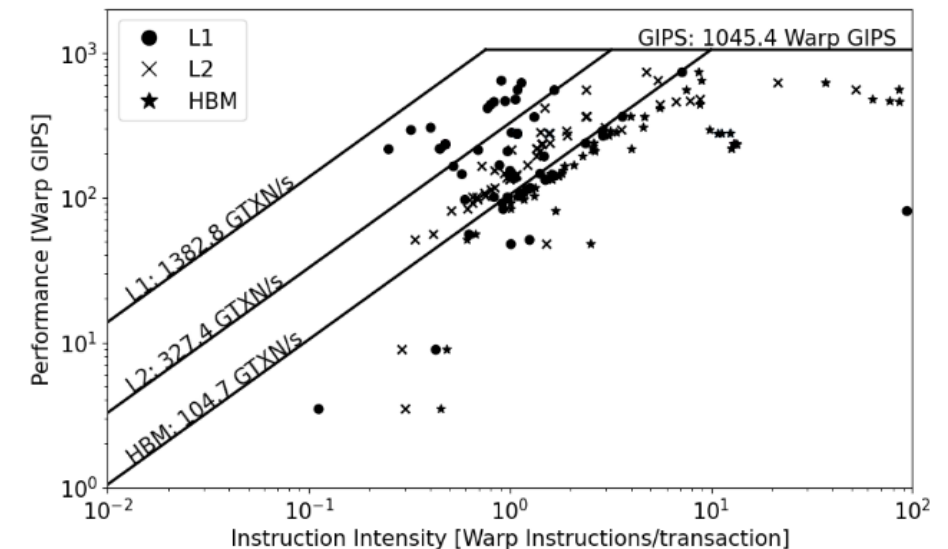
Use hardware metrics to evaluate similarity (runtime sensitive to problem size)

Intel's Top-down analysis [1] Intel Sapphire Rapids CPU



- Core bound – stalls from execution starvation
- Memory bound – stalls related to the memory subsystem
- Fetch Latency – corresponds to instruction cache misses
- Fetch Bandwidth – inefficiency in the instruction decoders

Roofline analysis [2] NVIDIA H100 GPU



- L1 Transaction Rate = $L1_{\text{Transactions, total}} / \text{Time}_{\text{GPU}}$
- L2 Transaction Rate = $L2_{\text{Transactions, total}} / \text{Time}_{\text{GPU}}$
- HBM Transaction Rate = $HBM_{\text{Transactions, total}} / \text{Time}_{\text{GPU}}$
- Performance (IPS) = $\text{Instructions}_{\text{warp}} / \text{Time}_{\text{GPU}}$

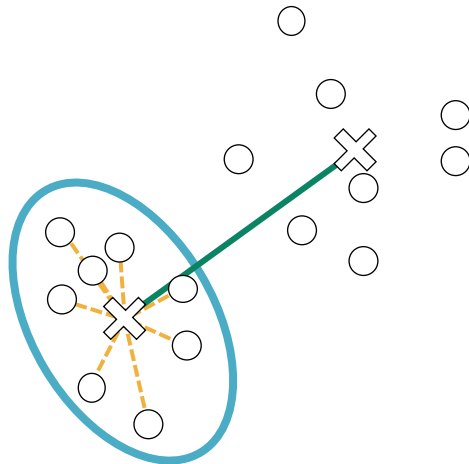
Does the proxy represent computational kernels from the application?

$$d(p, q) = \sqrt{(p_{corebound} - q_{corebound})^2 + \dots}$$



Euclidean distance $d(p, q) = \sqrt{(p_1 - q_1)^2 + (p_2 - q_2)^2 + \dots + (p_n - q_n)^2}$

where p_i is the value of the i -th metric in the hardware metrics list for kernel p



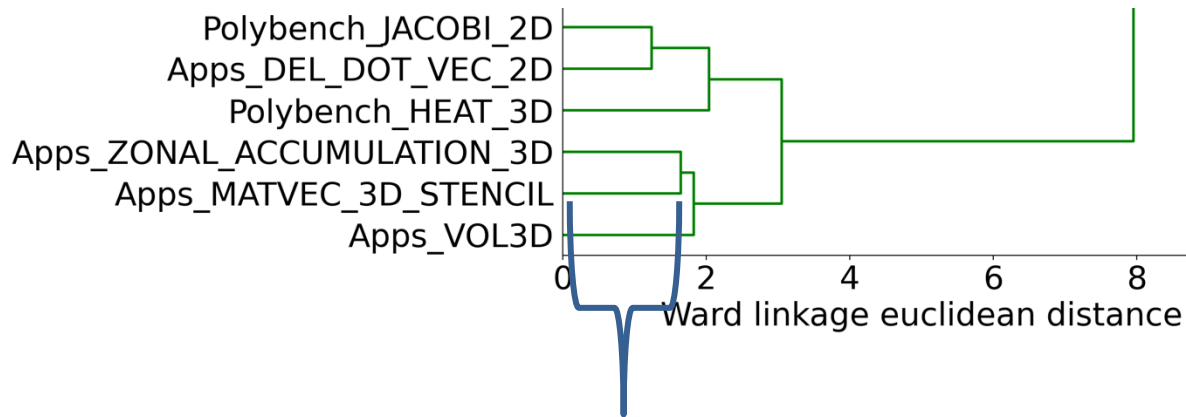
1. **Partitioning** – are kernels in the “correct” cluster
2. **Compactness** – The average *Euclidean distance* of each kernel to the cluster centroid
3. **Separation** – The average *Euclidean distance* between cluster centroids

Quantify similarity using partitioning, compactness, and separation

We use clustering to group kernels with similar performance

Agglomerative (ward linkage)

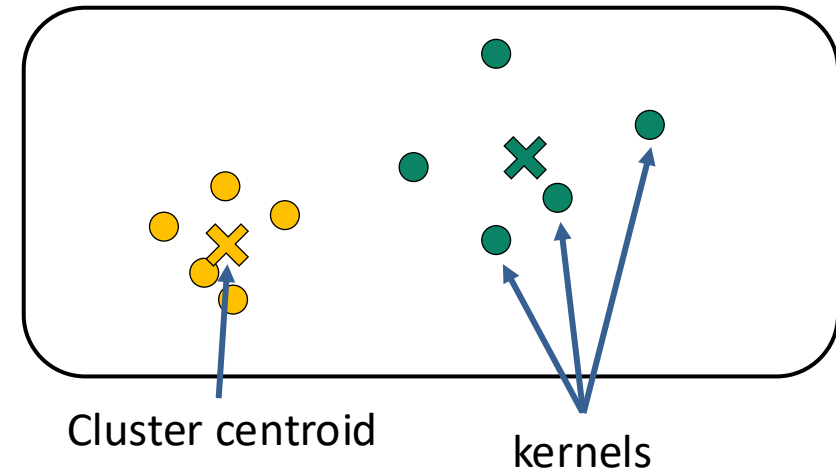
Create clusters “bottom-up” minimizing within-cluster variance



Distance at which “clusters” are merged. The first merge level indicates the two most similar kernels.

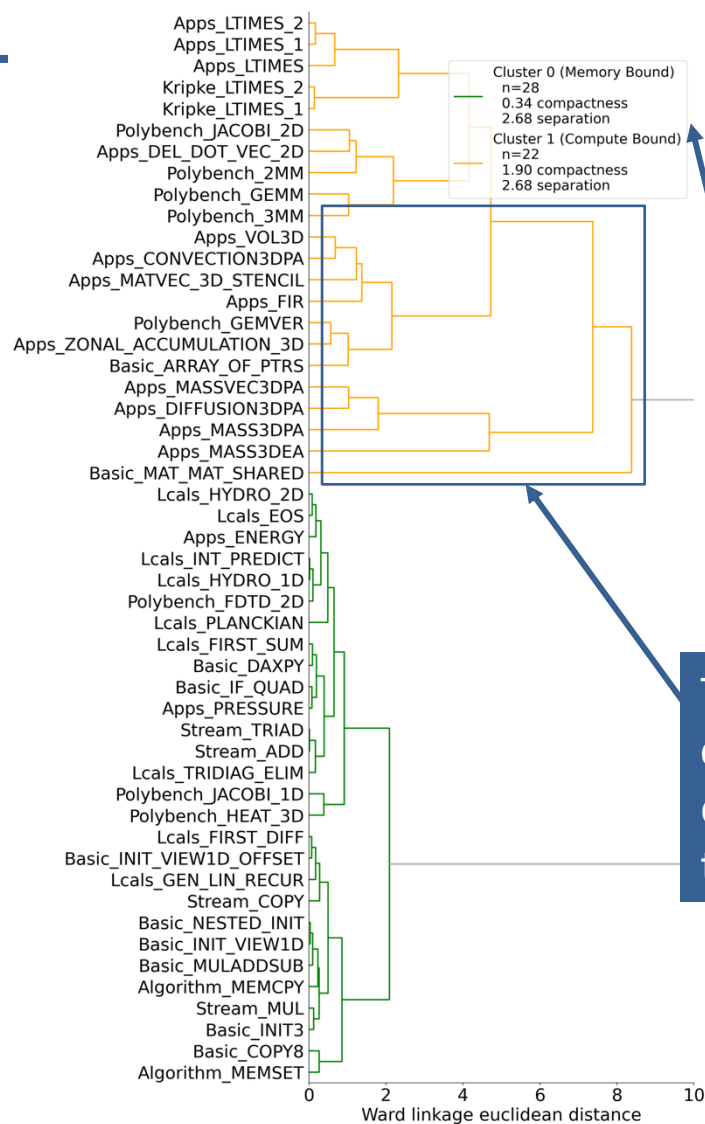
K-means (Lloyd’s algorithm)

Create clusters iteratively updating centroids to minimize within-cluster variance



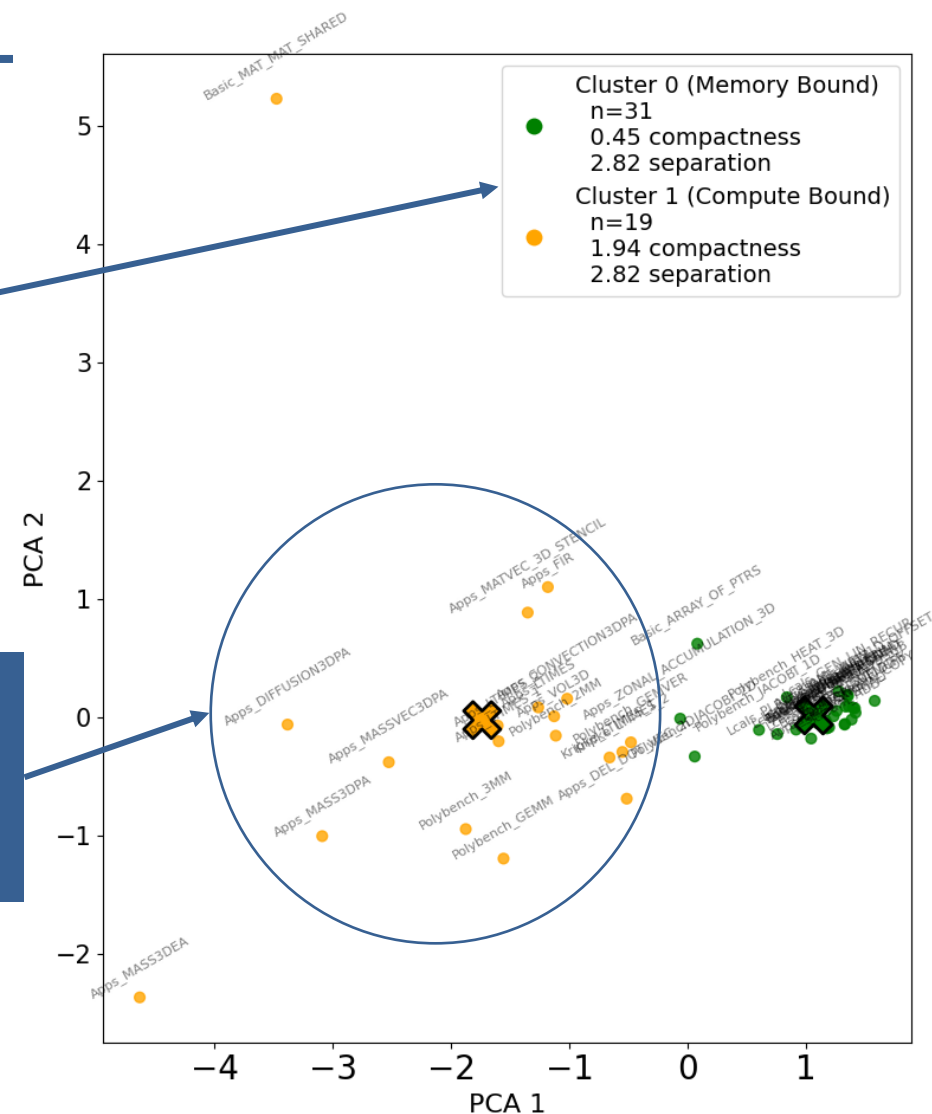
Do different clustering algorithms group kernels the same way?

Performance similarity on CPUs: Memory bound vs. Compute (core) bound

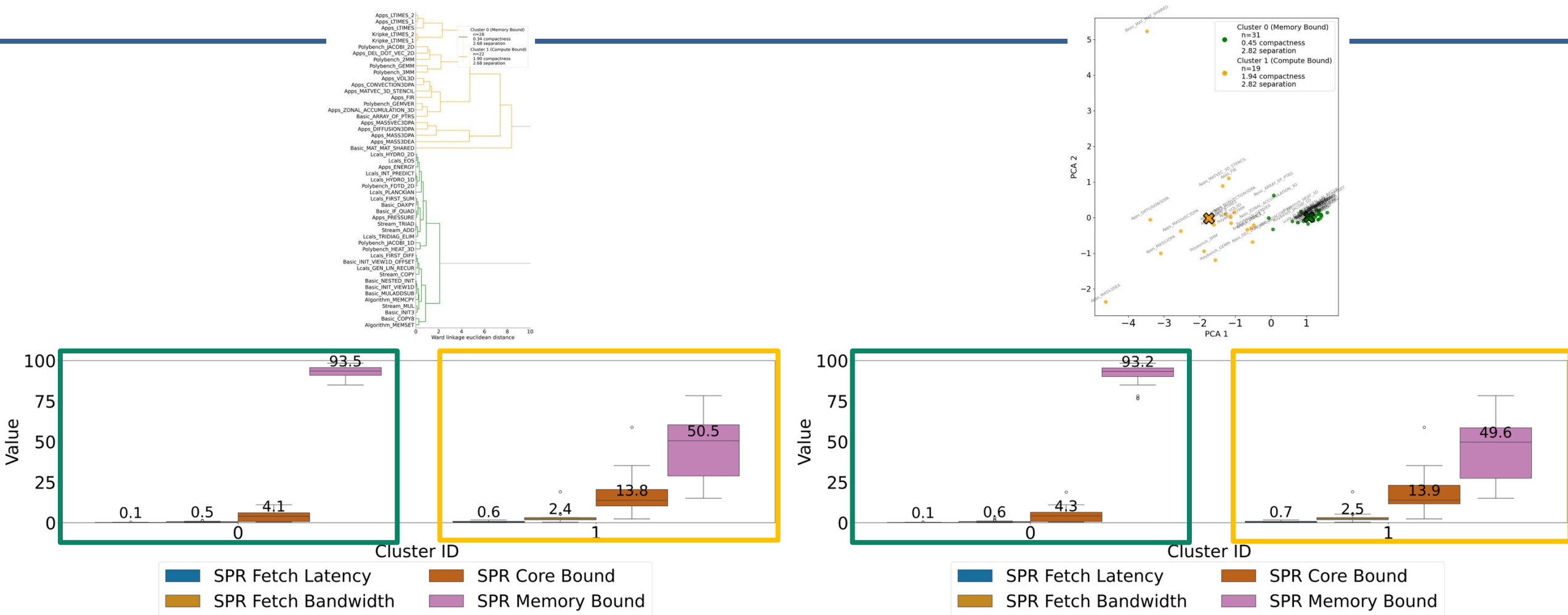


Partitioning of kernels is different for each clustering method, indicating that this subset of metrics alone may not be sufficient to categorize all the kernels.

The compute bound cluster is **not compact**, indicating the kernels in this cluster are not as similar as the kernels in the memory bound cluster.

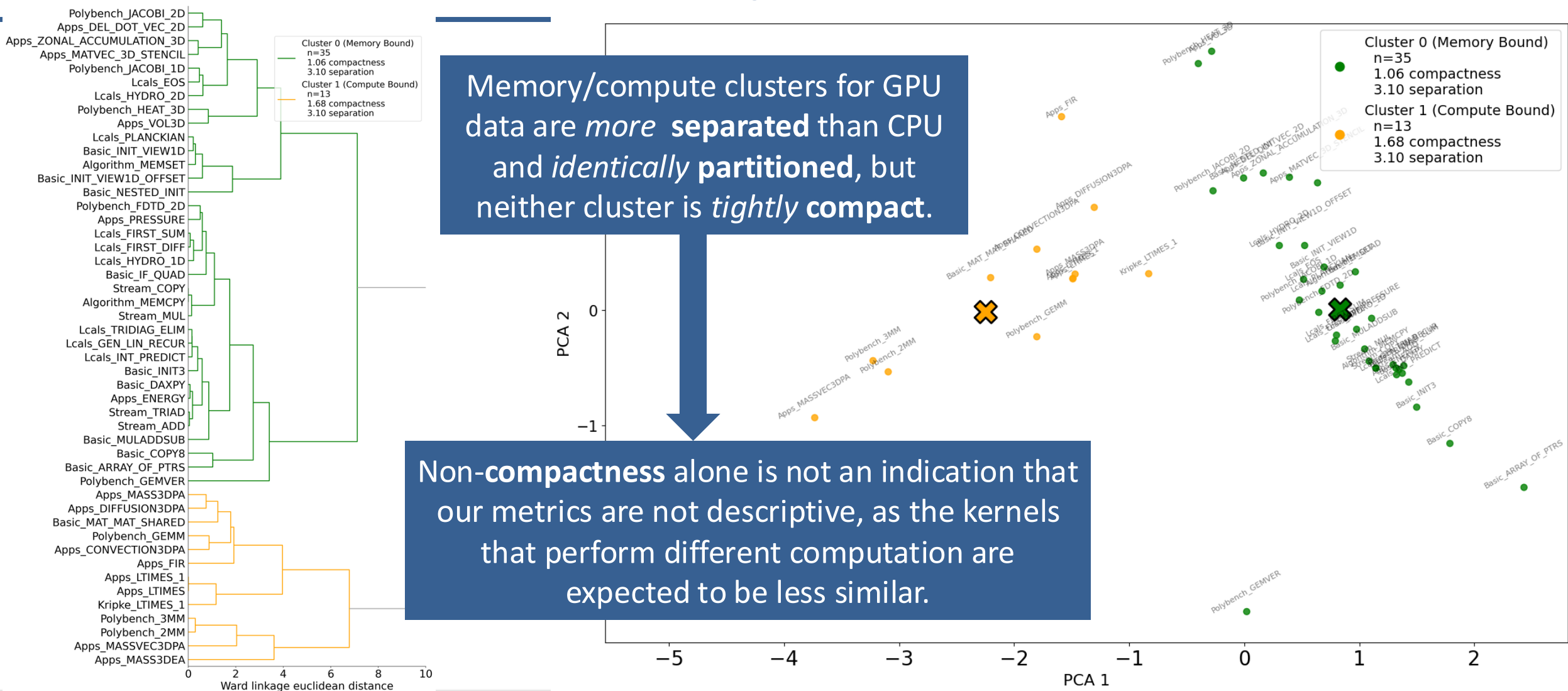


Performance similarity on CPUs: Memory bound vs. Compute (core) bound



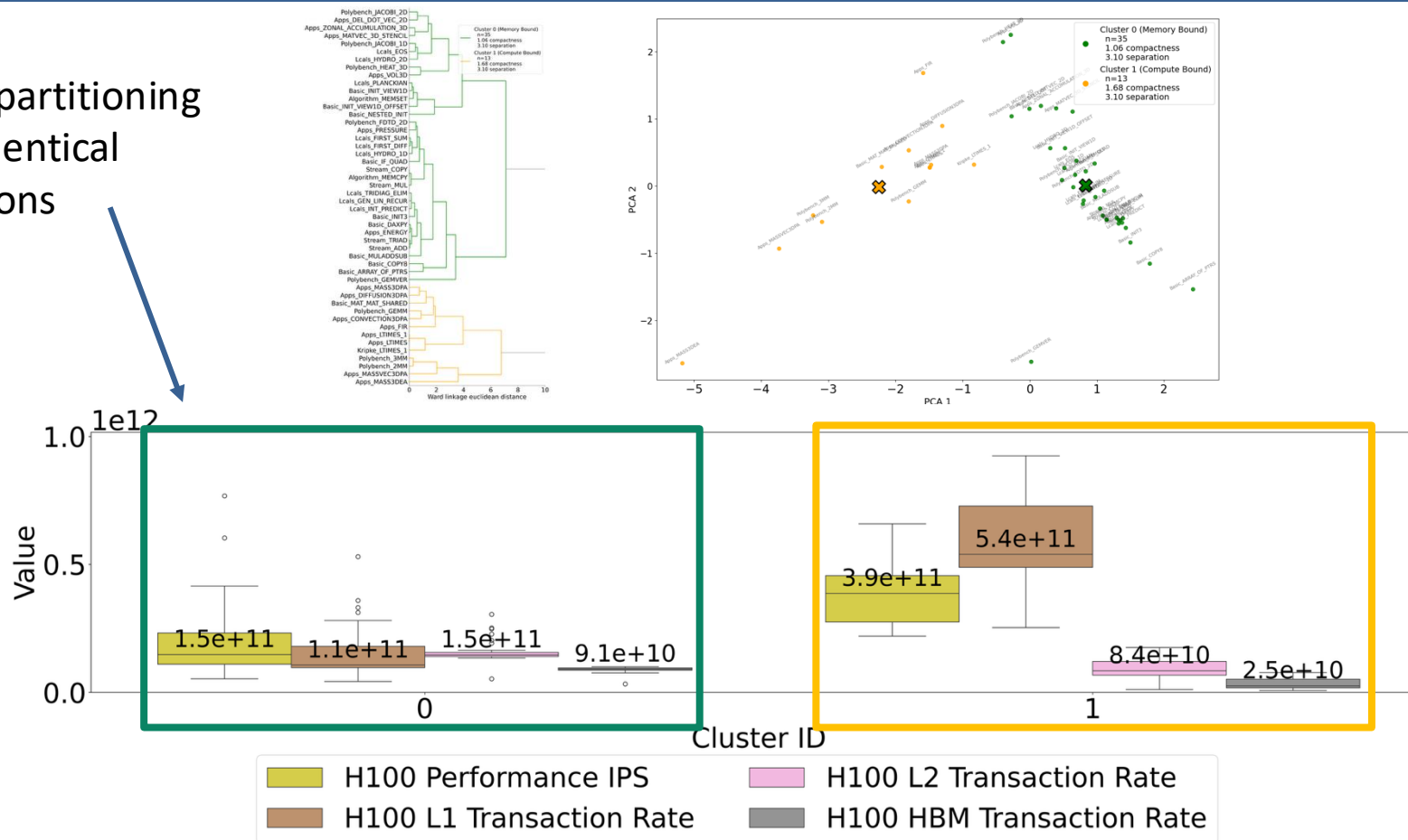
Memory bound kernels characterized by a high memory bound metric value (>90%). Compute bound kernels have significantly higher relative core bound values (~3x) along with higher fetch bandwidth/latency (~4-7x).

Performance similarity on GPU – memory (L2/HBM transaction rate) bound kernels versus compute (IPS) bound



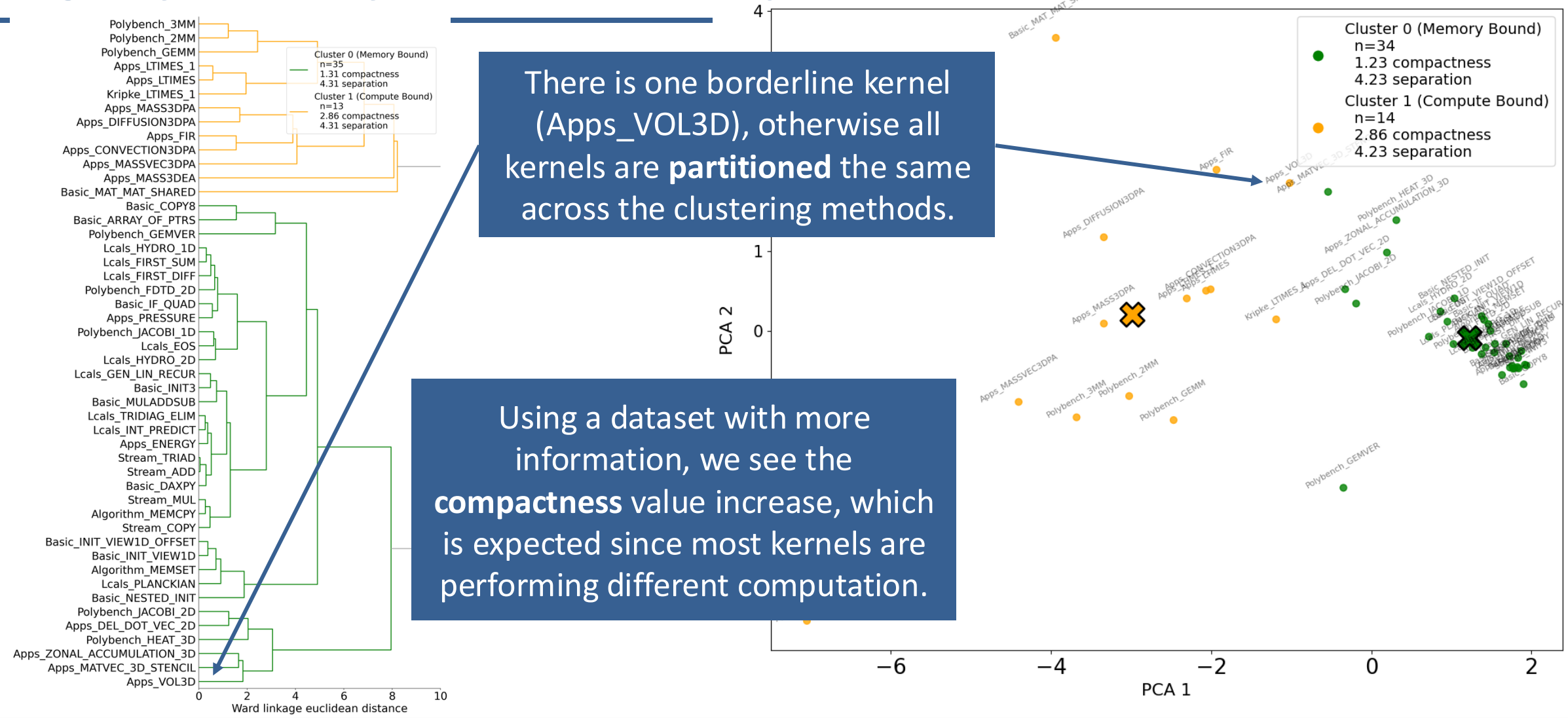
Performance similarity on GPU – memory (L2/HBM transaction rate) bound kernels versus compute (IPS) bound

identical partitioning
implies identical
distributions



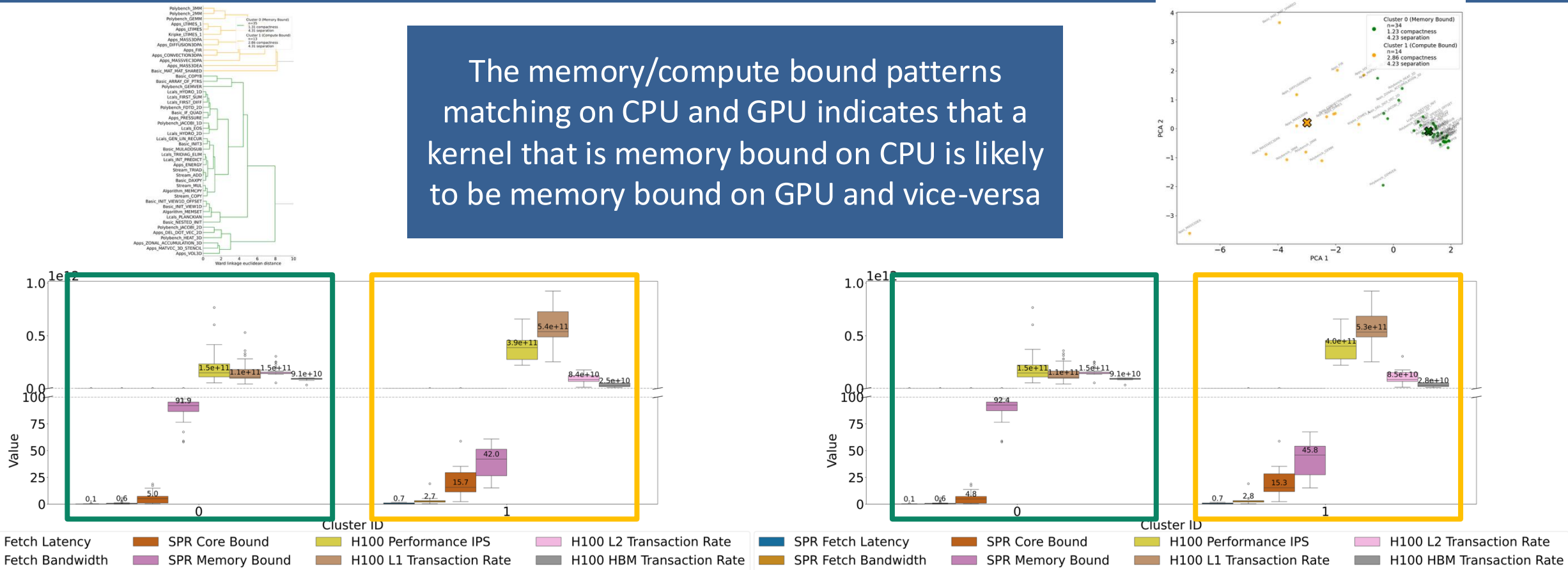
Memory bound kernels characterized by high L2 and HBM transaction rates (2-3x), where the compute bound kernels have high IPS and L1 transaction rates (2-5x).

Multi-platform performance similarity – CPU and GPU patterns align by memory bound and compute bound



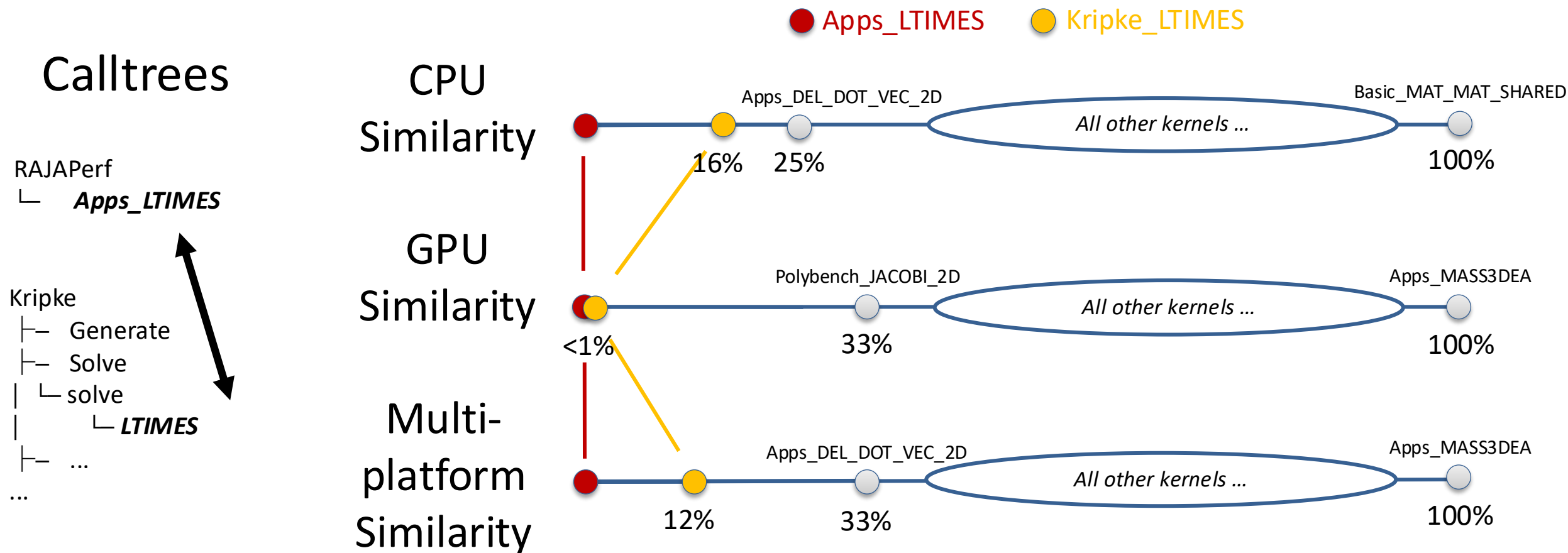
Multi-platform performance similarity – CPU and GPU patterns align by memory bound and compute bound

The memory/compute bound patterns matching on CPU and GPU indicates that a kernel that is memory bound on CPU is likely to be memory bound on GPU and vice-versa



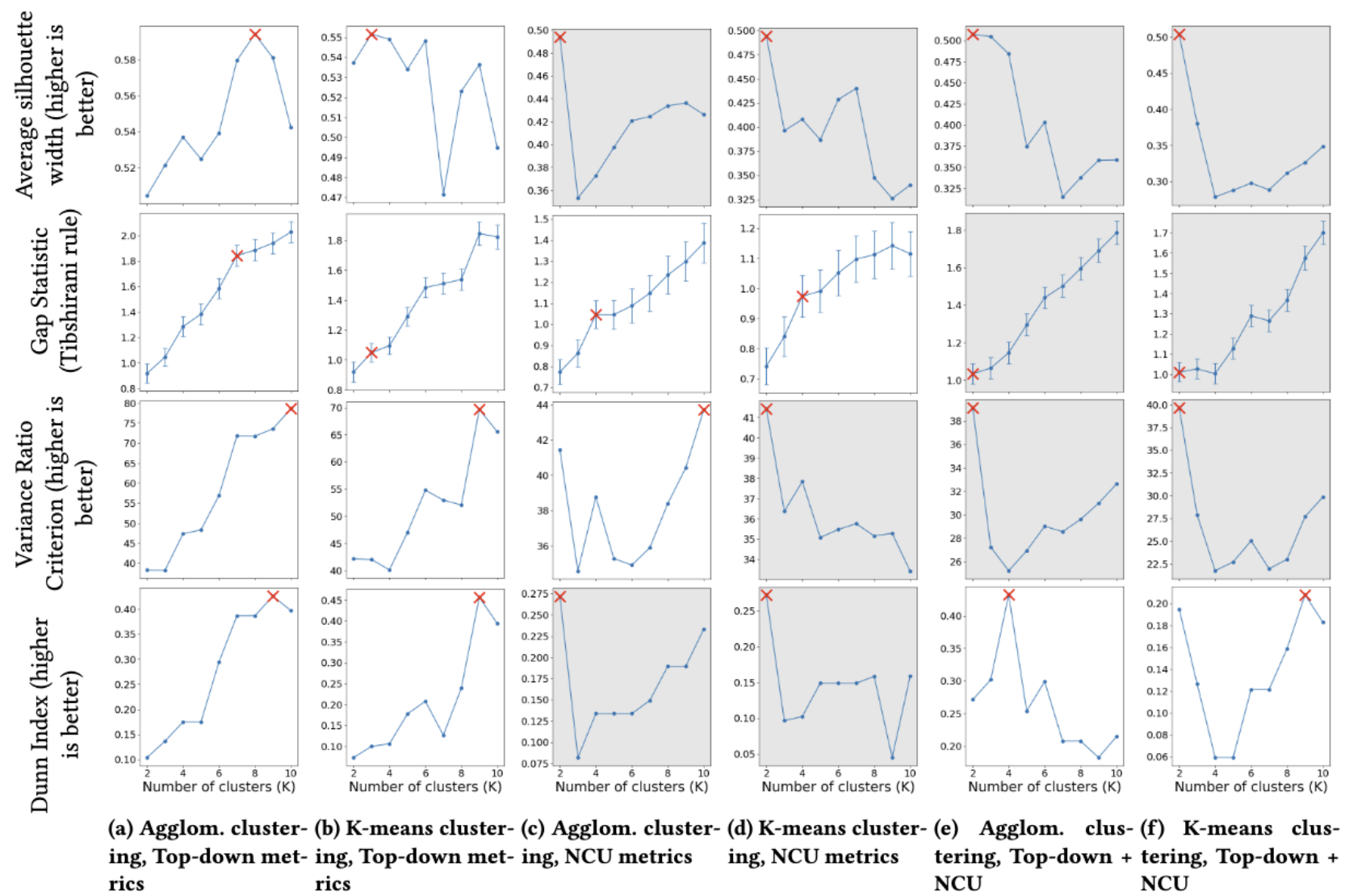
Memory bound kernels characterized by high memory bound, L2 and HBM transaction rates, where the compute bound kernels have high core bound, IPS and L1 transaction rates.

Similarity of multiple implementations of the LTIMES kernel



Kripke LTIMES is consistently the most similar kernel to RAJAPerf LTIMES

Selecting the number of clusters



Takeaways

- It is possible to classify kernels into memory and compute bound categories across multiple architectures with a handful of performance metrics.
- We can define performance similarity metrics to evaluate the similarity of many kernels across different applications.
- We can leverage performance data to evaluate whether a given kernel in a benchmark is representative of a kernel in an application.



CASC

Center for Applied
Scientific Computing



**Lawrence Livermore
National Laboratory**

Disclaimer

This document was prepared as an account of work sponsored by an agency of the United States government. Neither the United States government nor Lawrence Livermore National Security, LLC, nor any of their employees makes any warranty, expressed or implied, or assumes any legal liability or responsibility for the accuracy, completeness, or usefulness of any information, apparatus, product, or process disclosed, or represents that its use would not infringe privately owned rights. Reference herein to any specific commercial product, process, or service by trade name, trademark, manufacturer, or otherwise does not necessarily constitute or imply its endorsement, recommendation, or favoring by the United States government or Lawrence Livermore National Security, LLC. The views and opinions of authors expressed herein do not necessarily state or reflect those of the United States government or Lawrence Livermore National Security, LLC, and shall not be used for advertising or product endorsement purposes.