

Onur Cankur¹, Brian Austin², Abhinav Bhatele¹

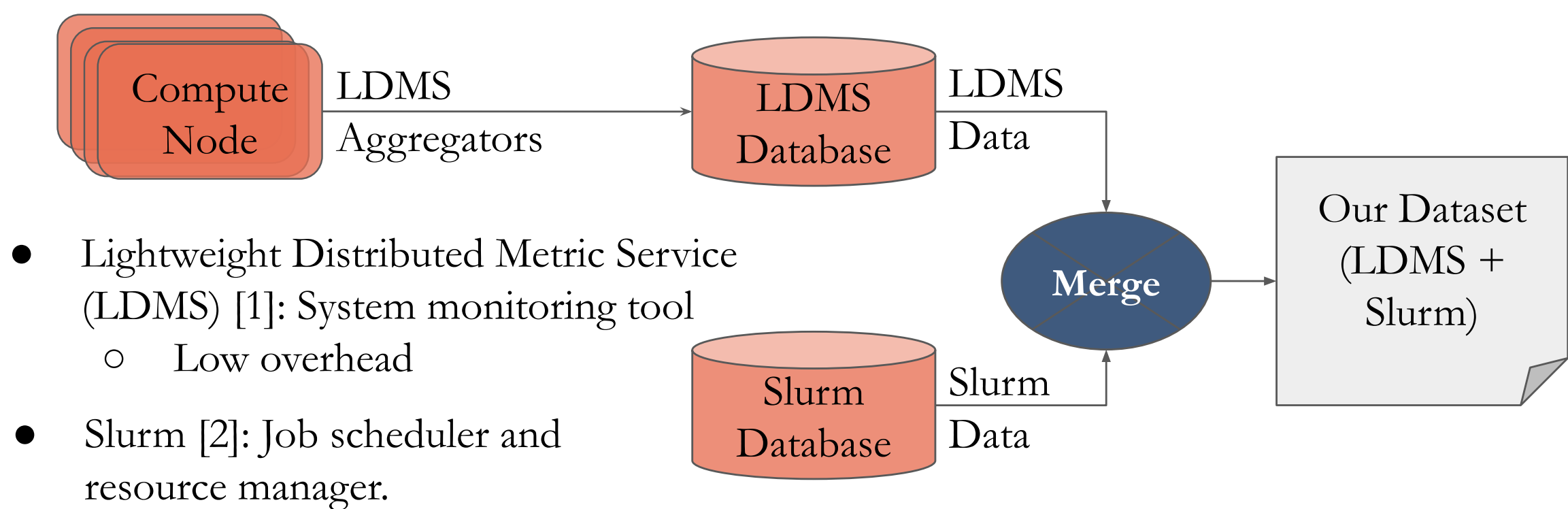
¹Department of Computer Science, University of Maryland

²NERSC, Lawrence Berkeley National Laboratory

Abstract

GPGPU-based clusters and supercomputers have grown significantly in popularity over the past decade. While numerous GPGPU hardware counters are available to users, their potential for workload characterization remains underexplored. In this work, we analyze previously overlooked GPU hardware counters collected via the Lightweight Distributed Metric Service (LDMS) on Perlmutter. We examine spatial imbalance, defined as uneven GPU usage within the same job, and perform a temporal analysis of how counter values change during execution. Using temporal imbalance, we capture deviations from average usage over time. Our findings reveal inefficiencies and imbalances that can guide workload optimization and inform future HPC system design.

Data Gathering and Preparation



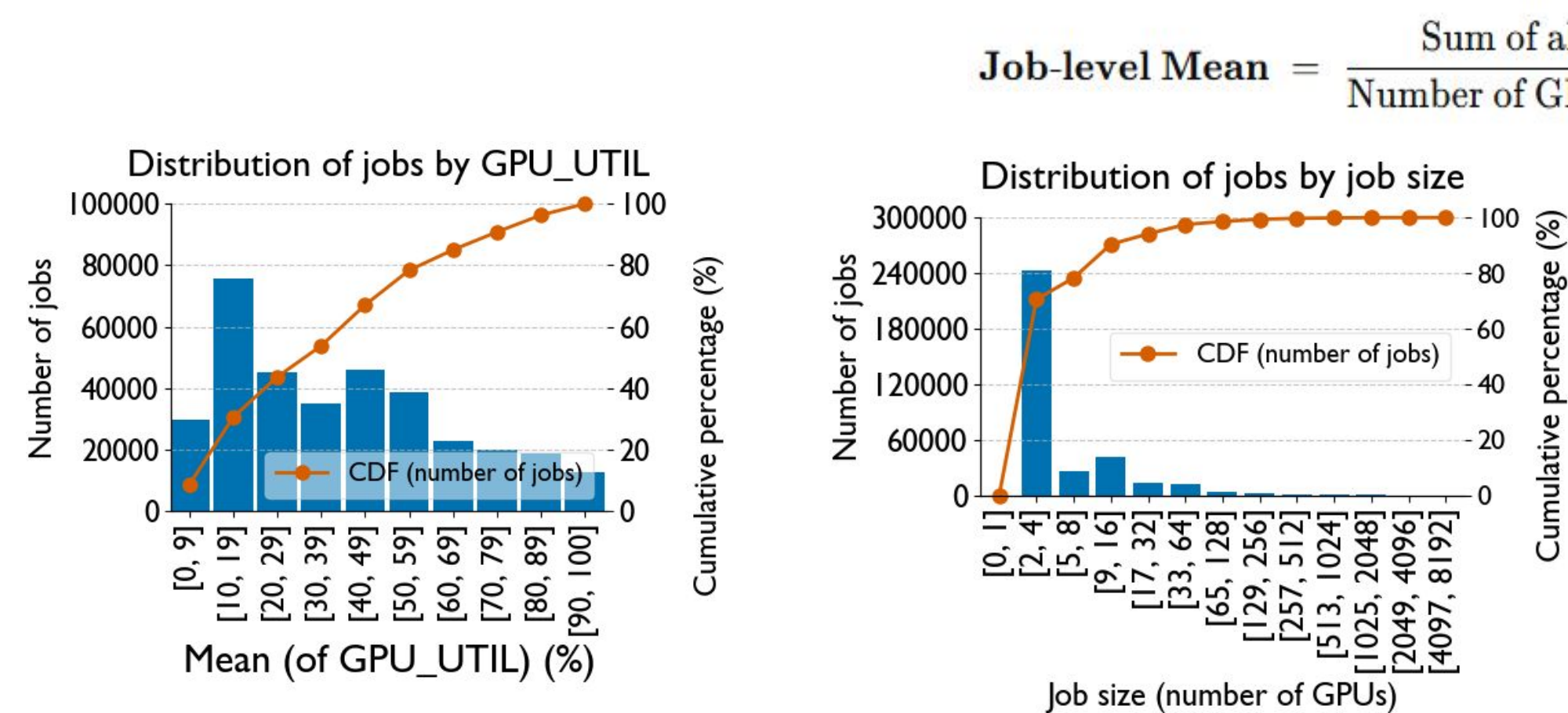
Monitoring Data Used

- We retrieved performance counter measurements collected by using the Data Center GPU Manager (DCGM) [3], configured with a 10-second sampling rate.
- We retrieved ~4 months of data, spanning August 16 to December 13, 2023.
- Our dataset includes information about 345,154 jobs after cleaning and preprocessing.

Counter Name	Short Name	Description
DCGM_FL_DEV_GPU_UTIL	GPU_UTIL	The fraction of time during which at least one kernel was executing on the GPU.
DCGM_FL_PROF_PIPE_FP{16/32/64/TENSR}_ACTV_ACTIVE	FP{16/32/64/TENSR}_ACTV	The fraction of cycles the FP{16/32/64/Tensor} cores were active.

Overview of Data

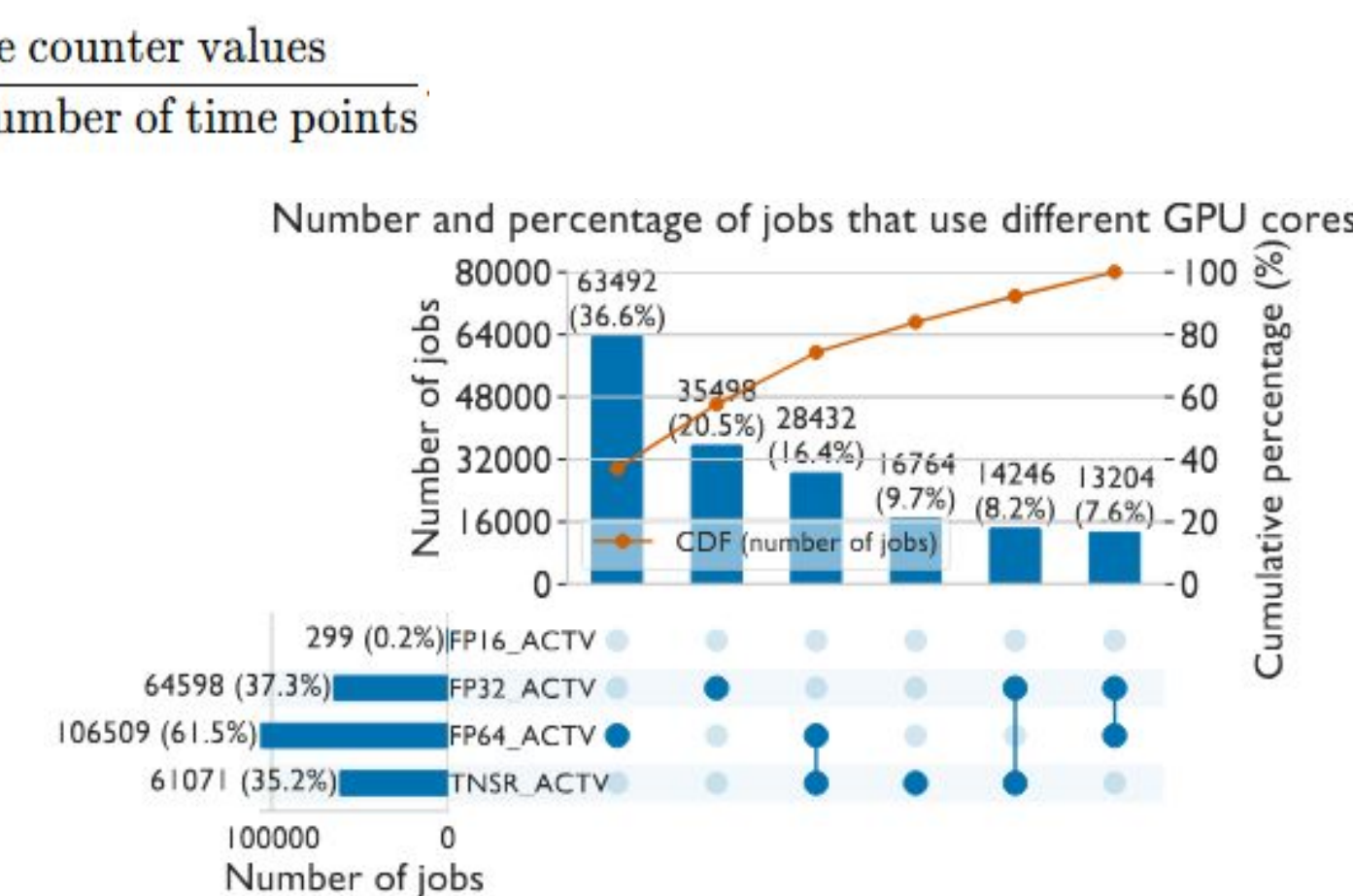
Question 1: How does overall GPU utilization and job count vary with job size (number of GPUs allocated)?



43% of jobs run at a mean GPU_UTIL of less than 30%.

Most jobs use only a single GPU node.

Question 2: Are there differences in how GPU jobs utilize FP16, FP32, FP64, and Tensor cores?

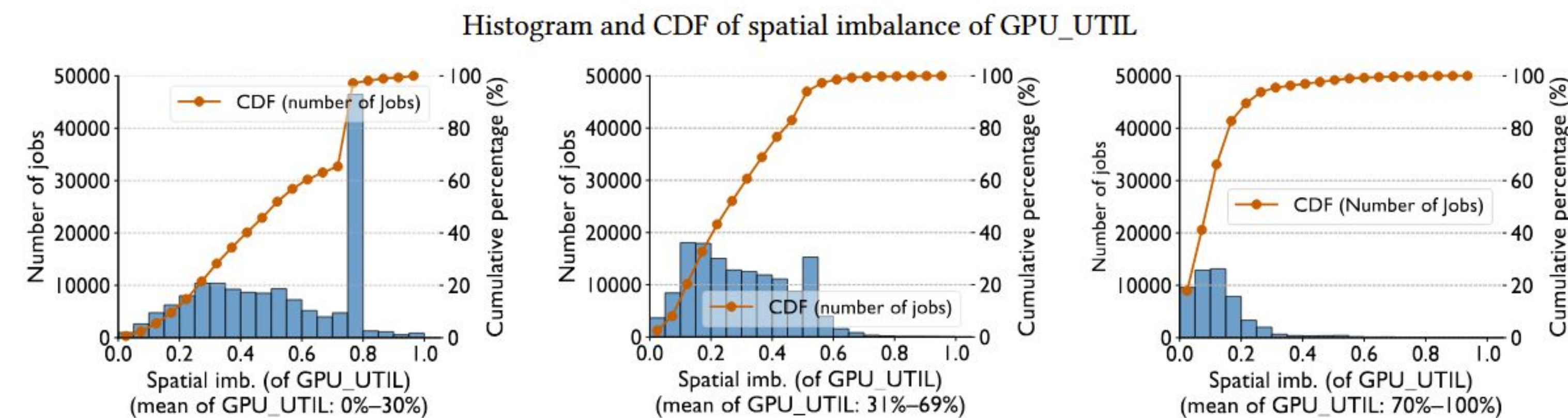


Over 60% of jobs use FP64 exclusively. Tensor operations co-occur with both FP32 and FP64.

Analyzing the Spatial Behavior of Jobs

Question 3: Do jobs that allocate multiple GPUs use them evenly, or are some GPUs heavily loaded while others stay idle?

$SI(job, window) = 1 - \frac{\text{Sum of counter values for all GPUs in the time window}}{\text{Highest per-GPU total in that window} \times \text{Number of GPUs}}$ Measures how evenly a job spreads work across its GPUs.



The jobs that have higher overall usage have less spatial imbalance.

Acknowledgements

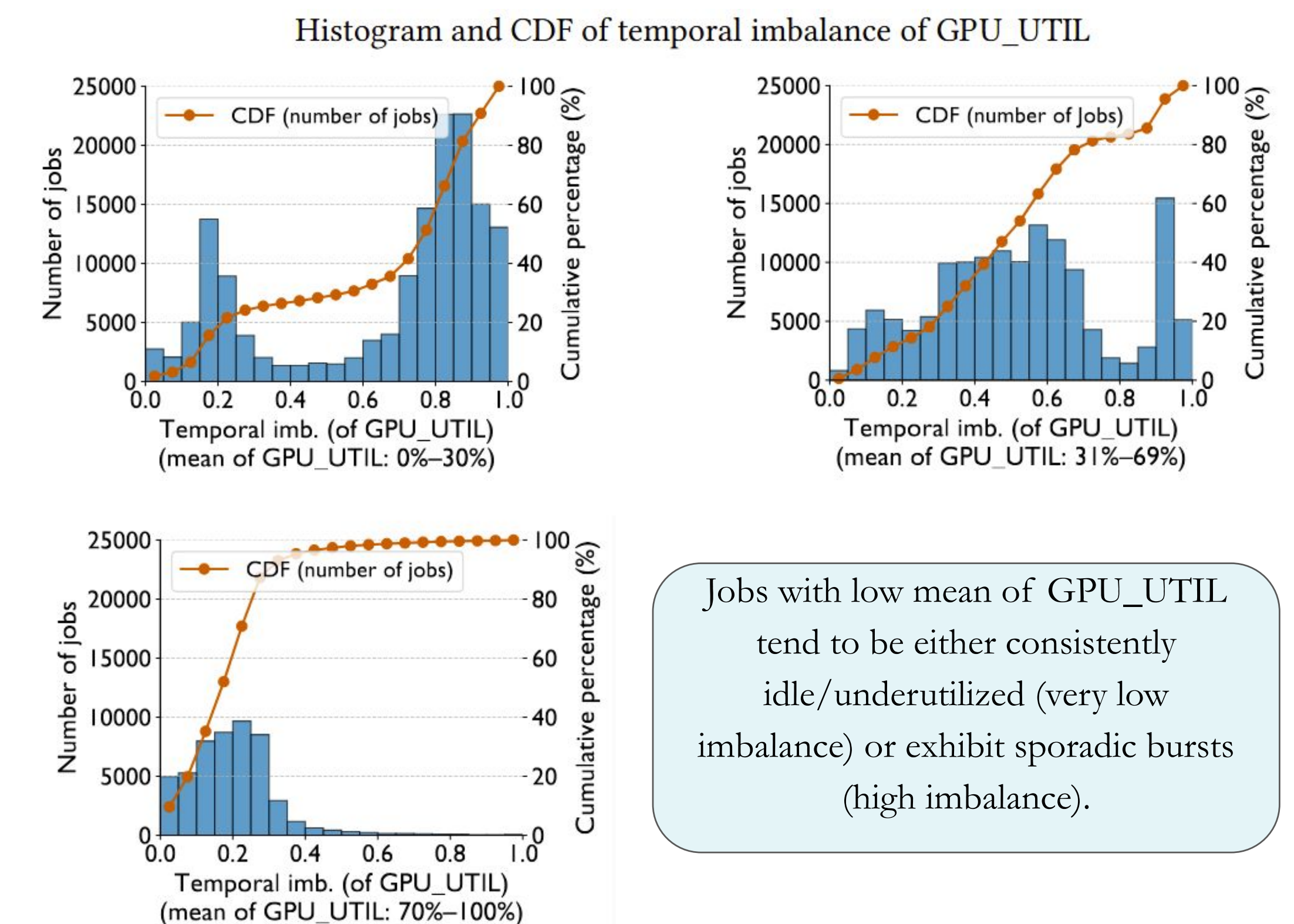
This material is based upon work supported by the National Science Foundation under Grant No. 2047120. This research used resources of the National Energy Research Scientific Computing Center (NERSC), a U.S. Department of Energy Office of Science User Facility, operated under Contract No. DE-AC02-05CH11231 using NERSC award DDR-ERCAP0034262.

Analyzing the Temporal Behavior of Jobs

Question 4: Are GPUs used consistently over time within a single job, or do we see bursts of heavy usage and idle periods?

$TI(job, gpu) = 1 - \frac{\text{Sum of counter values of the GPU over time}}{\text{Peak counter value} \times \text{Number of time points}}$

Quantifies how evenly a hardware counter is used during a job's runtime.



Discussion and Conclusion

- Uneven workload distribution across GPUs suggests potential for requesting fewer GPUs or improving load balancing.
- Perlmutter shows high demand for double precision (scientific/engineering) and lower demand for ML/AI jobs.
- Job-level computation phases mainly drive inconsistent usage patterns over time.
- Jobs with high GPU hours but low mean GPU_UTIL are prime optimization targets, identified by our methods.

References

- [1] Aggelastos, Anthony, et al. "The lightweight distributed metric service: a scalable infrastructure for continuous monitoring of large scale computing systems and applications." SC14: Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis. IEEE, 2014.
- [2] 2020. Slurm Workload Manager. <https://slurm.schedmd.com/documentation.html>
- [3] NVIDIA Corporation. 2023. NVIDIA Data Center GPU Manager (DCGM). <https://developer.nvidia.com/dcgmm>. Accessed: 2024-10-15.

