

Diagnosing Parallel I/O Bottlenecks in HPC Applications

Peter Harrington
Baskin School of Engineering
University of California Santa Cruz
Email: pharring@ucsc.edu

Abstract—High-Performance Computing (HPC) applications are generating increasingly large volumes of data (up to hundreds of TBs), which need to be stored in parallel to be scalable. Parallel I/O is a significant bottleneck in HPC applications, and is especially challenging in Adaptive Mesh Refinement (AMR) applications because the structure of output files changes dynamically during runtime. Data-intensive AMR applications run on the Cori supercomputer show variable and often poor I/O performance, but diagnosing the root cause remains challenging. Here we analyze logs from multiple levels of Cori’s parallel I/O subsystems, and find bottlenecks during file metadata operations and during the writing of file contents that reduced I/O bandwidth by up to 40x. Such bottlenecks seemed to be system-dependent and not the application’s fault. Increasing the granularity of file-system performance data will help provide conclusive causal relationships between file-system servers and metadata bottlenecks.

I. INTRODUCTION

High-performance computing systems have undergone great advancements in scale and capacity over the previous decades, which has brought about a significant increase in the complexity and parallelism of supercomputer architectures. In such HPC systems, achieving peak parallelism and performance in I/O will be an important step to building any feasible exascale computing system [1]. Identifying parallel I/O bottlenecks and their cause will improve efficiency and scalability at both the application level as well as the system level.

An HPC system uses multiple layers of software and hardware to implement parallel I/O. System logs exist for various levels of this stack, but comprehensive analysis of I/O performance is challenging due to the size, granularity, and varying sources of performance data.

II. METHODS

This work integrates performance logs from two key levels of Cori’s parallel I/O subsystems: application-level POSIX and Message Passing Interface (MPI) I/O traces, coming from the Darshan library [2], and performance data logged by the Lustre Monitoring Tool (LMT) [3], a monitoring tool for the Object Storage Targets (OSTs) and Metadata Servers (MDSs) of the Lustre parallel file system. We focus on the I/O behavior of two sample HPC AMR applications, *HyperCLaw* [4] and *Chombo* [5].

The I/O performance logs from large, highly parallel applications with are large and unwieldy to process. To efficiently analyze all of this information, Apache Spark [6] was used to

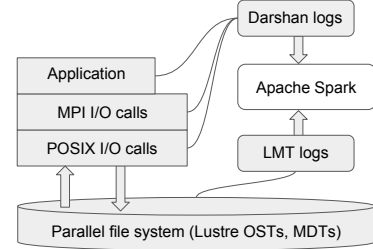


Fig. 1. A diagram of the analysis workflow. Darshan logs collect application-level I/O traces, while the LMT logs collect file-system usage data. Logs are parsed and analyzed in parallel using Apache Spark.

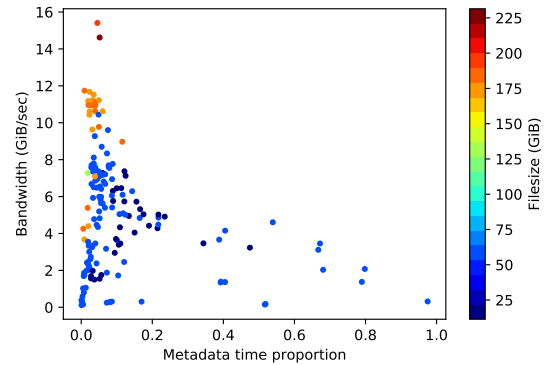


Fig. 2. The write bandwidth relative to the proportion of total I/O time spent in metadata operations, for all 79 runs of *HyperCLaw*. Bottlenecks happened when metadata operations dominated total I/O time (bottom right) and when writing of file contents dominated total I/O time (bottom left). Some outlying cases existed where bottlenecks in both components led to a roughly equal contribution to total I/O time from each (bottom center).

accelerate log-parsing and make I/O activity for specific times and MPI ranks easily accessible. A schematic diagram of this workflow is given in Fig. 1.

III. EXPERIMENTAL RESULTS

Even with optimized file system stripe settings, some runs showed significant slowdowns that caused as much as a 40x reduction in I/O bandwidth. These occurred when metadata operations or data write operations, and sometimes both, took much longer than usual and dominated the I/O time of the application. The relationship between write bandwidth

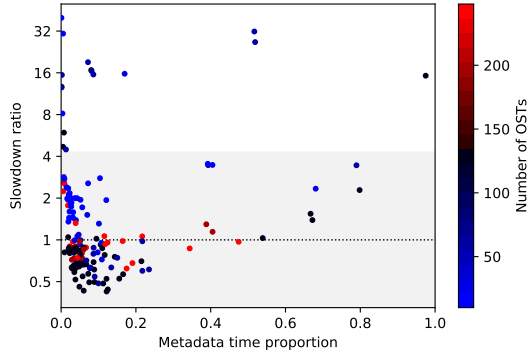


Fig. 3. The slowdown ratio of each run’s actual bandwidth relative to normal bandwidth (log scale), with the 90% range shown in gray. Besides runs with too few (blue) or too many (red) OSTs, most runs with a 2x or more slowdown had file system bottlenecks.

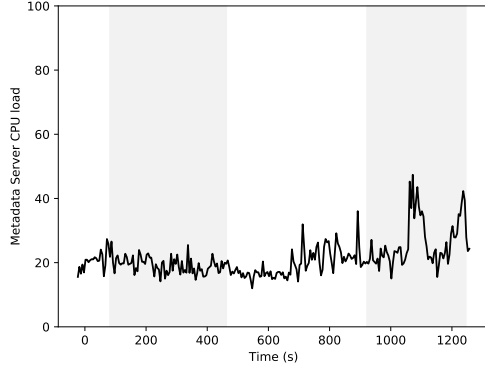


Fig. 4. The CPU load of the primary metadata server for the two files of a run where both metadata activities and data writing were bottlenecked. The write period of each file is shaded in gray.

and amount of I/O time spent in metadata operations can be seen in Fig. 2. Figure 3 shows the slowdown ratio of each run with respect to the normal bandwidth of runs with similar parameters. Other runs with identical settings to the bottlenecked cases were observed to run perfectly fine, so it appears that slowdowns during metadata operations and data writing were not the fault of the application.

File-system logs sampled activity every 5 seconds, and MDS CPU load was only measured for the primary MDS but not the four others. This limited granularity obscured the relationship between file system traffic and application-side I/O performance in bottlenecked runs. For example, even in the worst-performing run that saw a bottleneck in both metadata and data write operations for both files, only one of the file write periods coincided with elevated file system traffic (see Fig. 4).

We also observed an additional bottleneck in the applications with HDF5 collective writes, showing several trailing writes occurring after the main write output phase (see Fig. 5). These writes took up to 5.5 seconds, despite writing less than 5

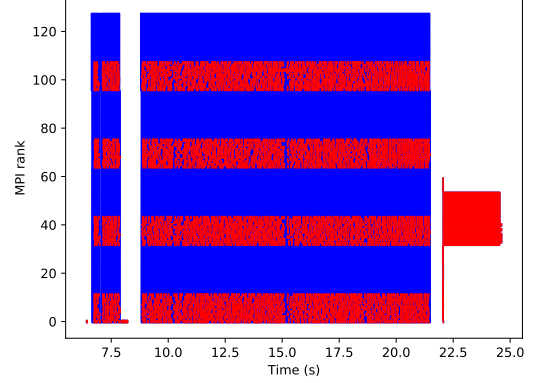


Fig. 5. MPI (blue) and POSIX (red) writes for each of the 128 MPI processes during a 104 GB Chombo run. Trailing writes occur near 22 seconds.

kB of data (a small fraction of the total output). Their duration increased when the total file size and number of processes increased. Such behavior is inefficient and does not scale.

IV. CONCLUSIONS

While the largest slowdowns in these application runs occurred during data writing, there were also significant slowdowns during metadata operations that seemed not to be the application’s fault. Most runs (90%) spent within 23% of total I/O time in metadata operations, and anything higher resulted in degraded performance. As traffic on HPC file systems continues to rise, it is likely that parallel I/O bottlenecks arising during metadata operations will become more frequent and severe.

Repeating this analysis with higher-resolution performance data for file-system activities should enable pinpointing such metadata bottlenecks to individual servers. During runs that did not experience any bottlenecks, many periods of metadata or data write operations took between 0.5-5 seconds, so knowing the load on file-system servers over shorter intervals is essential to finding correlations between file-system traffic and application-specific I/O performance.

REFERENCES

- [1] Kogge, P. et al. *ExaScale Computing Study: Technology Challenges in Achieving Exascale Systems*. Univ. of Notre Dame, CSE Dept. Tech. Report TR-2008-13. (2008)
- [2] Carns, P., et al. “Understanding and Improving Computational Science Storage Access through Continuous Characterization”. *Trans. Storage*, 7, 3. (2011)
- [3] Wartens, H. C. M. & Garlick, J. LMT - The Lustre Monitoring Tool. <https://github.com/chaos/lmt/wiki/>. Developed at Lawrence Livermore National Lab. (2010)
- [4] Welcome, M., et al. “Performance Characteristics of an Adaptive Mesh Refinement Calculation on Scalar and Vector Platforms”. *Proceedings of the 3rd Conference on Computing Frontiers*. (2006)
- [5] Adams, M., et al. “Chombo Software Package for AMR Applications - Design Document”. Lawrence Berkeley National Laboratory Technical Report LBNL-6616E. (2015)
- [6] Zaharia, M., et al. “Spark: Cluster Computing with Working Sets”. *Proceedings of the 2nd USENIX Conference on Hot Topics in Cloud Computing*. (2010)