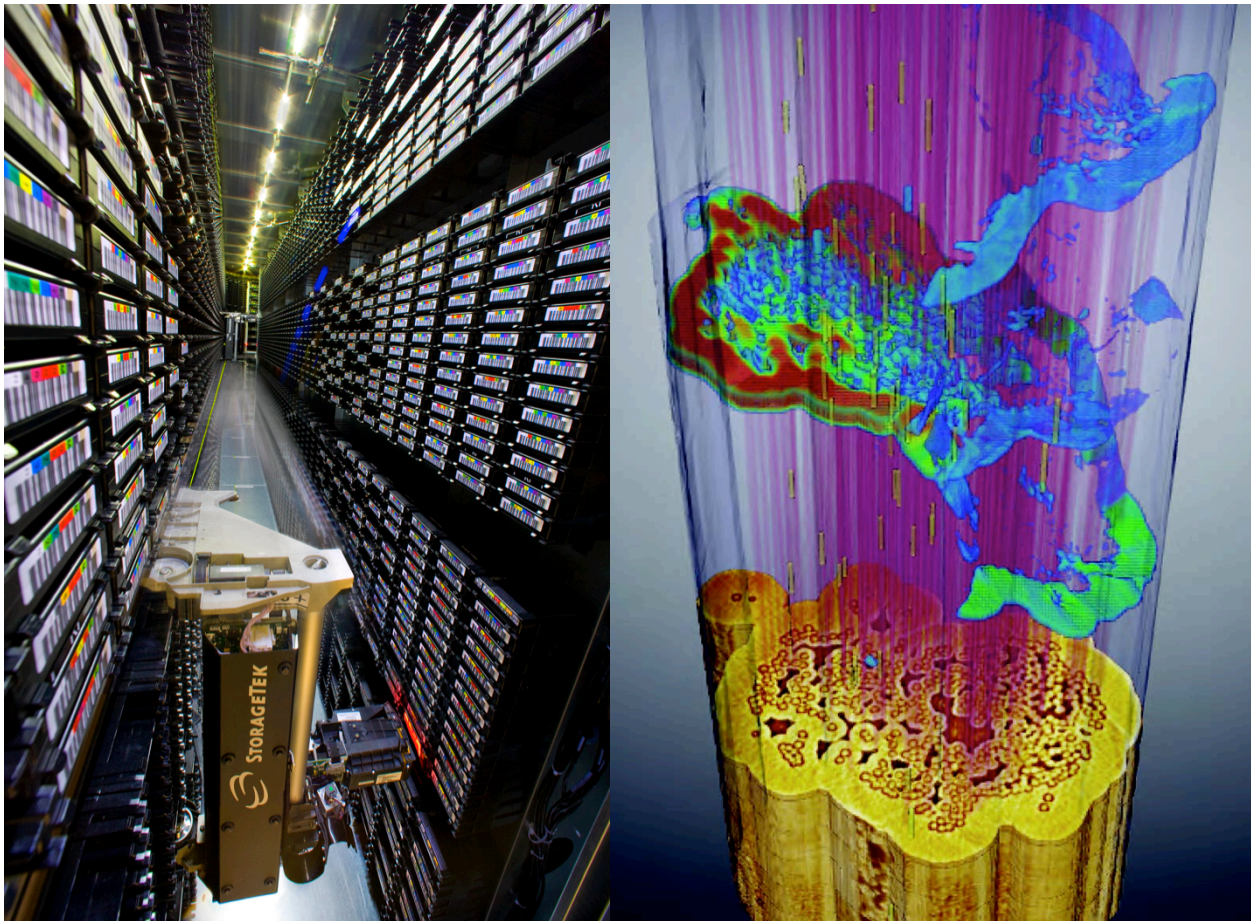


DOE High Performance Computing Operational Review (HPCOR)

Enabling Data-Driven Scientific Discovery at HPC Facilities



June 18-19, 2014
Oakland, CA



Editors

Richard Gerber, National Energy Research Scientific Computing Center, Lawrence Berkeley National Laboratory

William Allcock, Argonne Leadership Computing Facility, Argonne National Laboratory

Chris Beggio, Advanced Simulation and Computing, Sandia National Laboratories

Stuart Campbell, Spallation Neutron Source, Oak Ridge National Laboratory

Andrew Cherry, Argonne Leadership Computing Facility, Argonne National Laboratory

Shreyas Cholia, National Energy Research Scientific Computing Center, Lawrence Berkeley National Laboratory

Eli Dart, Energy Sciences Network, Lawrence Berkeley National Laboratory

Clay England, Oak Ridge Leadership Computing Facility, Oak Ridge National Laboratory

Tim Fahey, Advanced Simulation and Computing, Lawrence Livermore National Laboratory

Fernanda Foerster, Oak Ridge Leadership Computing Facility, Oak Ridge National Laboratory

Robin Goldstone, Advanced Simulation and Computing, Lawrence Livermore National Laboratory

Kevin Harms, Argonne Leadership Computing Facility, Argonne National Laboratory

Jason Hick, National Energy Research Scientific Computing Center, Lawrence Berkeley National Laboratory

David Karelitz, Advanced Simulation and Computing, Sandia National Laboratories

Laura Monroe, Advanced Simulation and Computing, Los Alamos National Laboratory

Prabhat, National Energy Research Scientific Computing Center, Lawrence Berkeley National Laboratory

David Skinner, National Energy Research Scientific Computing Center, Lawrence Berkeley National Laboratory

Julia White, Oak Ridge Leadership Computing Facility, Oak Ridge National Laboratory

September 17, 2014

On the cover: Left: Tape libraries of the HPSS archival storage system at the National Energy Research Scientific Computing Center (NERSC). Right: Three-dimensional rendering from computed microtomography data showing matrix cracks and individual fiber breaks in a ceramic matrix composite specimen tested at 1,750 C. Each of numerous ceramic samples is imaged with powerful X-ray scattering techniques over time to track crack propagation and sample damage, producing prodigious amounts of data. Increasingly, DOE imaging facilities are collaborating with DOE HPC centers to solve extreme data management and analysis challenges. Image credit: Hrishi Bale and Rob Ritchie, Lawrence Berkeley National Laboratory

Table of Contents

INTRODUCTION	5
EXECUTIVE SUMMARY	7
KEY FINDINGS.....	7
OPPORTUNITIES FOR COLLABORATIONS AMONG DOE HPC CENTERS.....	8
REPORTS FROM BREAKOUT SESSIONS	9
SYSTEM CONFIGURATION FOR DATA ANALYTICS	9
VISUALIZATION/ <i>IN SITU</i> ANALYSIS.....	13
DATA MANAGEMENT POLICIES	20
SUPPORTING DATA-PRODUCING FACILITIES AND INSTRUMENTS	25
INFRASTRUCTURE	30
USER TRAINING FOR DATA-INTENSIVE SCIENCE.....	35
WORKFLOWS	40
DATA TRANSFER.....	47
ATTENDEES	53

DISCLAIMER

This report was prepared as an account of a workshop sponsored by the U.S. Department of Energy. Neither the United States Government nor any agency thereof, nor any of their employees or officers, makes any warranty, express 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 any agency thereof. The views and opinions of document authors expressed herein do not necessarily state or reflect those of the United States Government or any agency thereof. Copyrights to portions of this report (including graphics) are reserved by original copyright holders or their assignees, and are used by the Government's license and by permission. Requests to use any images must be made to the provider identified in the image credits.

Ernest Orlando Lawrence Berkeley National Laboratory is an equal opportunity employer.

Ernest Orlando Lawrence Berkeley National Laboratory
University of California
Berkeley, California 94720 U.S.A.



NERSC is located at the Lawrence Berkeley National Laboratory, which is operated by the University of California for the US Department of Energy under contract DE-AC02-05CH11231.

This work was supported by the Directors of the Office of Science, Office of Advanced Scientific Computing Research (ASCR) and the National Nuclear Security Administration, Advanced Simulation and Computing (ASC) Program.

This is LBNL report xxx published xxx.

Introduction

U.S. Department of Energy (DOE) High Performance Computing (HPC) facilities are on the verge of a paradigm shift in the way they deliver systems and services to science and engineering teams. Research projects are producing a wide variety of data at an unprecedented scale and level of complexity, with community-specific services that are part of the data collection and analysis workflow. The value and cost of data relative to computation is growing and, with it, a recognition that concerns such as reproducibility, provenance, curation, unique referencing, and future availability are going to become the rule rather than the exception in scientific communities. Addressing these concerns will impact every facet of facility operations and management. The optimal balance of hardware architectures may change. Greater emphasis will be given to designing software to optimize data movement relative to computational efficiency. Policies about what data is kept, how long it is kept, and how it is accessed will need to adapt. Data access for widespread scientific collaborations will become more important. Processes and policies that ensure proper and secure release of information will need to evolve to both maintain data protection requirements and meet future data sharing demands.

On June 18-19, 2014 representatives from six DOE HPC centers¹ met in Oakland, CA at the DOE High Performance Operational Review (HPCOR) to discuss how they can best provide facilities and services to enable large-scale, data-driven scientific discovery at the DOE national laboratories. Discussions were focused around separate topics in the following eight breakout sessions:

1. **System Configuration (Session D1SA):** What are the hardware characteristics of a good data analytics system, including compute and storage? What does it look like? Should an HPC system and a data system be the same or different? What percentage of resources should be allocated to compute vs. data/I-O? What storage technologies and tools are being used and which new ones are being considered? *Co-Chairs: Jason Hick, NERSC; Clay England, Oak Ridge National Laboratory*
2. **Visualization/In Situ Analysis (Session D1SB):** What is needed to support *in situ* analysis and visualization? From hardware, software, and support perspectives? What visualization facilities and capabilities do you support for both local and remote users? *Co-Chairs: Prabhat, NERSC/Berkeley Lab; David Karelitz, Sandia National Laboratory; Laura Monroe, Los Alamos National Laboratory*
3. **Data Management Policies (Session D1SC):** What facilities and policies are in place for data retention and access? What are the challenges and possible solutions? Will centers be

¹ The Argonne Leadership Computing Facility (ALCF) located at Argonne National Laboratory; the National Energy Research Scientific Computing Center (NERSC) located at Lawrence Berkeley National Laboratory; the Oak Ridge Leadership Computing Facility (OLCF) located at Oak Ridge National Laboratory; and the Advanced Simulation and Computing (ASC) facilities at Lawrence Livermore National Laboratory, Los Alamos National Laboratory, and Sandia National Laboratories.

part of scientists' Data Management Plans? If so, how? What standards for data repositories and archives are in place and which ones do you plan to support? How is access to the broader community provided? How do you balance storage costs with data retention and access policies? How is data management planned? *Co-Chairs: Julia White, Oak Ridge National Laboratory; Bill Allcock, Argonne National Laboratory*

4. **Supporting Data-Producing Facilities and Instruments (Session D1SD):** What is your center doing to support data and its analysis from light sources, accelerators, satellites, etc.? *Co-Chairs: David Skinner, NERSC; Stuart Campbell, Oak Ridge National Laboratory.*
5. **Infrastructure (Session D2SA, Room 2):** What supporting infrastructure is needed to enable data-driven science? Which of these play primary roles in supporting your center's data-driven science: networking between resources, shared or local disk, archival storage, science data gateways or portals, consulting, and databases. *Co-Chairs: Robin Goldstone, Lawrence Livermore National Laboratory; Chris Beggio, Sandia National Laboratories.*
6. **User Training (Session D2SB):** What are the methods for effective user training for data-driven science? Visualization. Tools. Algorithms. Workflows. I/O. Who and where are the experts who provide training? Are they in the technical systems groups, services group, vendors, users? *Co-Chairs: Fernanda Foertter, Oak Ridge National Laboratory; Tim Fahey, Lawrence Livermore National Laboratory.*
7. **Workflows (Session D2SC):** What is being used? What works well and what is missing? What infrastructure and support is required? *Co-Chairs: Shreyas Cholia, NERSC; Kevin Harms, Argonne National Laboratory*
8. **Data Transfer (Session D2SD):** What WAN access is in place and what is needed? How do you handle data transfers in/out of your facility today (e.g. do facility staff conduct transfers or do users, what hardware/software do users use)? What drives the need for networking and what will HPC centers need to do to accommodate that need? *Co-Chair: Eli Dart, ESnet; Andrew Cherry, Argonne National Laboratory*

In addition to the questions associated with the breakout session topics given above, attendees were also asked to consider the following:

- What are your major strategies and initiatives over the next 5-10 years? How do they affect staffing levels?
- What are your current efforts and/or site configuration in this area?
- What are your mandates and constraints?
- How do you forecast future needs and requirements?
- What are the biggest challenges and biggest gaps between what you can do today and what will be required in 5-10 years?
- What opportunities exist for productive collaborations among DOE HPC centers?
- Describe some best practices that you think are effective as well as lessons learned that would be helpful to other centers.

Executive Summary

To support the needs of the U.S. Department of Energy data-producing facilities – whether they are High Performance Computing Centers, experimental facilities, or other collaborative projects like astronomical sky surveys – will require DOE HPC centers to change the way they have traditionally operated. The breakout groups at this review identified a number of key challenges and issues. Highlights of those findings are presented here, and additional key points are listed in the sections of this report associated with each breakout session.

Key Findings

- A tighter integration between experimental facilities (within BER and BES in particular) and ASCR HPC centers would decrease time to publication and enable new science from the national user facilities. Computational facilities need to be fully engaged with the DOE experimental facilities' current and future plans for their data and compute needs.
- Workflows – sequences of coordinated compute and data-centric operations – are a fundamental aspect of data-intensive science, but their use on modern HPC architectures is still new and there is little standardization. Workflows require documentation, support infrastructure (e.g., databases, master control nodes, virtual environments), and staff to facilitate site integration.
- Many data analysis workflows need access to diverse node types (e.g. some with TBs of memory) and storage systems that perform well with random or unpredictable access patterns. Traditional large HPC systems have not been designed to accommodate these workflows, and DOE centers are evaluating whether or not a single system can accommodate these needs or if separate systems will be required.
- A complex hierarchy of storage resources will be required to satisfy data-intensive computing needs. Everything from fast local storage (probably solid-state drive, or SSD) to archival storage will be needed, and new software tools must be built to manage and monitor data flow through the different levels of storage.
- Current HPC job schedulers were not designed to handle high throughput computing or complex workflows. New schedulers are needed that are aware of storage and network resources and are able to handle millions of jobs.
- Current network security and data access policies pose significant challenges to data-intensive workflows. Data needs to flow seamlessly and at high performance from remote instruments to, and among, HPC centers and back to collaborators worldwide.
- Data management policies are nascent, and mandates to scientists and HPC centers are still developing. Centers need to be agile in such an environment. Issues such as data lifetimes, access, provenance, and curation must be considered.
- DOE and HPC centers need to develop performance and evaluation metrics that are relevant to data-intensive science and have less emphasis on computation.
- Data transfer nodes (DTNs) are key for moving data in and out of DOE HPC centers. To fully exploit increasing backbone bandwidth requires careful attention to DTN-to-file-system performance, local network topologies, tools performance, data set composition and other

factors.

- *In situ* and in transit visualization and analysis will become more common, and their design should be driven by the end user in collaboration with the other relevant parties (hardware, simulation, visualization, and analysis experts) to ensure success. The analysis routines must be first-class citizens with the same level of support (both developer and end user) as the simulation code.
- Users are facing a multitude of new challenges surrounding large data sets. They need training to learn about techniques and best practices for managing and analyzing data.

Opportunities for Collaborations Among DOE HPC Centers

- Create a set of relevant metrics and tools for measuring and reporting these metrics.
- Define a set of metrics to evaluate workflows, and offer recommendations based on various criteria such as performance, throughput, ability to handle different classes of problems, feature sets, ease of use, etc.
- Build a common infrastructure to accommodate multi-site workflows.
- Research the best ways to design and implement burst buffers.
- Create a set of benchmarks, or proxy applications, that are representative of common visualization and analysis workloads, that cover *in situ*, in transit, and post-hoc use cases, for purposes that include, but are not limited to, system testing, performance tuning, and optimization.
- Create software collaboration efforts like SDAV in general, and the VTK-m collaboration between the VisIt and ParaView teams in particular.
- Collaborate at all levels to support data-producing facilities, from the PIs to the experimental facility staff to HPC center staff.
- Investigate available options and alternatives to existing storage technologies and methods.
- Organize a “Birds of a Feather” (BoF) or workshop as part of a major conference such as Supercomputing for, for example, workflows or security in a shared data environment.
- Build a catalog of workflow tools along with their capabilities to provide sites and users with guidance in terms of what tools to choose for their jobs.
- Agree upon, provide, and regularly test a common tool set and use consistent tuning practices. Sites should develop a set of common best practices and policies.
- Share best practices to improve interoperability between HPC sites. Regular meeting of working groups consisting of representatives from multiple labs is one way to encourage this.
- Develop a strategy for data sharing (especially in-place sharing). It is important that this need be addressed in a way that provides adequate safeguards for security, accountability, and control of potentially sensitive data.
- Develop standards across facilities and centers to enable synchronization of data formats, workflows, and environments in a secure and transparent manner.
- Build on existing inter-facility partnerships to produce generalized services for data-intensive science.
- Users of DOE HPC centers have many common needs for training on data management and analysis. The centers can (and have begun to) coordinate and collaborate on training events on topics of common interest to users of all the centers.

Reports from Breakout Sessions

The reports from the breakout sessions follow.

System Configuration for Data Analytics

Session D1SA

Authors

Clay England (OLCF), Jason Hick (NERSC)

Contributors

Andrew Cherry (ALCF), Cory Lueninghoener (LANL ASC), Curt Canada (LANL ASC), Robin Goldstone (LLNL ASC), Jim Silva (LLNL ASC), Nick Wright (NERSC), Chris Beggio (SNL ASC), Bob Balance (SNL ASC), Glenn Lockwood (SDSC), Eli Dart (ESnet)

Introduction

The use of high performance computing (HPC) for data analytics is an area of increasing interest. Traditionally, HPC facilities have been used to produce large-scale predictive simulations with subsequent post-processing data analysis for knowledge discovery. But with the advent of large simulation data sets (generated at large HPC facilities) and large experimental data sets (generated at large experimental user facilities), the need for HPC user facilities to prepare for the onset of data analytics requests is clear. To prepare for this need, we discuss below the expected architectural and software requirements of such systems, the need for changing operational priorities, and opportunities for HPC facilities to increase their cross-facility collaborations.

Key Points

Resources needed to support data analytics on HPC Centers

- Tools that will move data up and down the persistent storage hierarchy
- Nodes with large memory (TBs)
- Fast low-latency I/O (probably SSD) that integrates into the storage hierarchy (scratch disk, shared networked disk, archival storage)
- Parallel file systems and network speeds that can support a data-intensive workload
- Utilization metrics that have less emphasis on computation
- The ability to run data-intensive workflows that were not designed to run under standard HPC schedulers
- Better tools to monitor I/O activity, data, and network systems

Opportunities for collaboration among centers

- Workload analysis and creating a set of common measurement tools and metrics
- Building the infrastructure to accommodate multi-site workflows
- Burst buffer research
- Centralized reporting

Discussion

Because of the relative newness of data analytics, participants discussed a variety of solutions in use at their facilities for data analysis and analytics. Some had traditional stand-alone InfiniBand or Ethernet-based commodity clusters purposefully dedicated to running Hadoop or other similar MapReduce technologies for unstructured data sets from observational and experimental data collections. In these cases, a user may run “Hadoop on demand” instances (e.g. Spot Hadoop, MyHadoop) on HPC clusters, often without any involvement from an operations team. Other participants combined data analytics with data analysis to create a suite of systems and services supporting this broader definition of data analytics. These facilities believe that data analytics generally involves the use of several different resources (e.g. clusters, DTNs, portals, visualization, and data processing) in post-processing of large-scale simulation data and is less dependent on a cluster running Hadoop/MapReduce.

While there is still disagreement as to what data analytics/analysis is, there is some consensus as to the specifications of a system that is capable of either task. In general, there are similarities between data analytics systems and traditional HPC systems. Data analytics application developers generally prefer a standard and homogeneous environment because much of the software development occurs on standard PC hardware. More specialized (BlueGene) or heterogeneous (accelerators) HPC systems present challenges to current data analytics software developers. Furthermore, most data analytics software has not been developed with an eye toward HPC system capabilities.

On the hardware side, workshop participants identified that the standalone Hadoop clusters generally have commodity storage components (e.g., each nodes’ local disks) that are prone to failure. This is in contrast to most HPC systems, which present large parallel shared file systems as their primary storage.

Future Architecture Requirements for Data Analytics Systems

The most common architectural requirement shared between data analytics applications is the need for a large amount of memory per core (or per node). This is true today (e.g., visualization clusters require large memory for the purpose of rendering) and will continue to remain true. Furthermore, fast, low-latency I/O will continue to be a common request for these machines. Local SSDs will be useful for fast, local storage; for example, currently SSDs are used for staging data, holding an HDF5 file, or as a local file system that is closer to the compute cores compared to a large, shared parallel file system. This node local storage should work in conjunction with the persistent storage hierarchy in use today for HPC workloads (e.g., shared parallel file systems).

GPUs (or other accelerators) may also be useful in data analytics workloads. Some Hadoop versions can make use of them, as can most analysis and visualization workloads.

Future Software Needs for Data Analytics Systems

The most pressing software need for data analytics is a tool that will move data up and down the persistent storage hierarchy (memory, node local persistent storage, parallel file system) to (1) get the data close to the core when it's needed and (2) aid in post processing once the data has been collected. For example, data movement is a large problem in the MapReduce workload where the node local file system might not be high capacity, requiring that data be migrated to the parallel file system as part of the job conclusion process. The use of hierarchical storage management (HSM) software could help migrate and purge data throughout the storage hierarchy. Other software needed to facilitate data analytics on HPC systems are search tools, R/SAS for statistical work, Hadoop, and noSQL databases.

Operational Priorities for Data Analytics Systems

Operational priorities and existing operational practices may need to change for systems dedicated to data analytics. For example, the use of schedulers (e.g. Moab, LSF, Slurm) is a common practice at almost all HPC facilities. Some of the current Hadoop-on-demand systems are scheduled and run alongside other HPC workloads under a standard HPC scheduling system. However, existing data analytics systems at different HPC centers are operated differently – some are scheduled while others are not, mostly depending on the users needs and experience.

Utilization metrics need to be redefined for data analytics workloads. Data analytics is a more read-heavy operation than most traditional HPC workloads. For example, some MapReduce workloads have been optimized for running on an HPC system; while the number of read operations was increased by four times, the overall time to solution was improved by a factor of 20. The big data sets used for the data analytics workloads coupled with the different ways in which these workloads make use of the system will influence how we charge (and account) for their usage. CPU core hours will become less important in data analytics workloads, while moving large data sets in and out will be the greatest cost in this area.

Today, operationally, we think of HPC centers in terms of peak Flop/s. With the shift toward a data-intensive workload, the typical breakdown of compute versus I/O and storage will likely be different. Determining the appropriate ratio common to all centers is likely not useful because different facilities have different compute and analysis needs. However, the order in which system hardware is chosen may change to:

1. Determine the memory/core needed for workloads
2. Determine the amount of SSD or persistent storage needed
3. Determine the parallel file system and network speeds needed for data-intensive computing
4. Allocate the remainder of the budget to Flop/s (CPUs, accelerators, many-core chips)

The above choices will be affected by the workload in the facilities. Because of the increased focus

on data and storage, there will be a need for better tools to monitor performance of data/storage and I/O on HPC systems.

Opportunities for Collaborations between HPC Centers

There are many opportunities for collaboration between HPC centers in data analytics and analysis. Workload analysis is one area where the centers could work together for data analytics workloads. Each center could identify their data-intensive workloads and share them with one another. This could lead to more workload analysis tools in common use at the centers (e.g., Darshan). Another area would address the increasing demand for multi-facility workflows. This increase is due partly to experimental facilities that don't want or can't replicate what HPC facilities have or can provide. On a related note, it was suggested that HPC centers explore storage and network quality of service (QoS), both locally and between facilities. There are some collaborations in place using flash with a parallel file system that would also help in furthering data analytics support. NERSC, SNL ASC, and LANL ASC are working together with Cray to develop a burst buffer solution to improve burst I/O capabilities for supercomputers; this should also provide great benefit to the data analytics and analysis community. We discussed establishing and publishing performance expectations for data ingest or export to demonstrate to scientists what performance to expect from facilities; this would suggest centralized reporting from multiple facilities to be helpful to users.

Future Challenges

Participants identified two main hurdles for data analytics in HPC environments. The main issue is the lack of software development that is needed for HPC centers to support data analytics. Becoming more involved in the development process or educating the developers to HPC needs/wants can alleviate this. The second large hurdle is multi-facility workflows, especially to address the issue of different security domains.

Visualization/*In Situ* Analysis

Session D1SB

Authors

Laura Monroe (LANL ASC), Prabhat (NERSC), David Karelitz (SNL ASC)

Contributors

Kevin Harms (ALCF), Bob Kares (LANL ASC), Ming Jiang (LLNL ASC), Jeff Long (LLNL ASC), Wes Bethel (LBNL), Doug Fuller (OLCF), Andy Wilson (SNL ASC), Kenneth Moreland (SNL ASC), Sean Ahern (OLCF)

Key Points

- *In situ* and in transit analysis and visualization will be a necessary component of modeling and simulation at exascale, but they will not completely replace post-processing capabilities
- In transit visualization and analysis could require one type of node for simulation and another for the visualization and analysis. This leads to issues of co-scheduling or job reservations containing multiple node types.
- Visualization and analysis capabilities integrated into the simulation code should be driven by the end user and result in a collaboration among the relevant parties (hardware, simulation, visualization and analysis, and the domain scientist) to ensure success.
- Since *in situ* visualization and analysis are integrated with the simulation code, there is a necessity for the *in situ* code to be a first-class citizen with the same level of support (both developer and end user) as the simulation code.
- For remote access, one of the biggest challenges is accessing resources behind firewalls and opening the necessary ports to enable remote users.
- For in transit and post processing, visualization and analysis nodes would ideally be a separate class of nodes from the compute nodes incorporating differing architectures and perhaps additional resources, such as extra memory or burst buffers. Regardless of node type, it is imperative that the visualization nodes have a high availability for interactive use, especially if data is being post-processed.

Opportunities for Collaborations Among HPC Centers

- Create a set of benchmarks, or proxy applications, that are representative of common visualization and analysis workloads, that cover *in situ*, in transit, and post-hoc use cases, for purposes that include, but are not limited to, system testing, performance tuning, and optimization.
- Participate in software collaboration efforts such as SDAV in general and the VTK-m collaboration between the VisIt and ParaView teams in particular.

Discussion

Overview

This breakout session was primarily devoted to discussions of current visualization and analysis best practices and the changes predicted to be necessary to support exascale platforms. The discussion touched on all aspects of the development cycle, from research and development to deployment and support, with issues and recommendations for all aspects present in the discussion below.

Traditional visualization and analysis uses a post-processing or post-hoc paradigm in which visualization and analysis output files are saved to disk and the visualization and analysis is performed after the simulation finishes running.

Since CPU and memory performance is scaling much faster than disk bandwidth, the post-hoc paradigm is expected to become more difficult to support on exascale machines. Two technologies that reduce the amount of data written to disk are *in situ* and in transit visualization and analysis. Both of these methodologies perform visualization and analysis while the simulation is running, with the main difference being where the visualization/analysis is performed.

In situ visualization/analysis is a tight coupling of the simulation code with the visualization and analysis code where the visualization/analysis runs on the same nodes, and potentially the same executable, as the simulation and does not transfer the simulation data over the mesh or network. In transit visualization/analysis is a looser coupling in which the simulation transfers data from the simulation compute nodes to a set of visualization/analysis nodes for processing. The set of in transit nodes could have a different node architecture, although this isn't necessary. However, the in transit concept does imply that the visualization process does not run on the same nodes as the simulation. *In situ* implies that the visualization is run on the same nodes as the simulation. Both *in situ* and in transit visualization run before analysis products are written to disk.

There was general agreement among the participants that post-processing visualization is an essential part of the scientific process and will not vanish with exascale. However, the scope of this session for this meeting was *in situ* and in transit technologies. We will therefore not discuss post-processing techniques for exascale in great detail in this document; however, it may be appropriate to revisit this issue for exascale planning purposes.

Topics Discussed

Post-Processing vs *In Situ* Visualization/Analysis

There was some discussion regarding the trade-offs between post-processing and *in situ* visualization and analysis. Post-processing enables interactive data exploration and discovery, while visualization and analysis performed during the simulation typically only saves data products that are known a priori. There will always be a tradeoff between these two paradigms, but there was consensus that both were necessary and likely required to be supported in the near future.

For some simulations there is a need to compare and store results over a period of years, which is much easier with a post-processing paradigm. For others, particularly uncertainty quantification (UQ) where thousands of parameterized runs may be performed, *in situ* is a better fit.

A combined approach may also be advantageous. The use of certain *in situ* analysis techniques might give enhanced data reduction and might also furnish techniques for narrowing down the data space. Either/or may not be the best approach to this question.

What is needed to support *in situ* analysis and visualization including hardware, software, and support perspectives?

The discussion on this topic centered on the need for both hardware and software vendor support and the need for a unified co-design effort among the hardware teams, simulation codes, visualization/analysis team members, and domain experts. There was strong consensus among the participants that early and sustained involvement of domain experts is needed to implement *in situ* and in-transit visualization successfully.

Supporting both *in situ* and post-hoc analysis/visualization on next-generation platforms will require a co-design effort as the visualization software is adapted to new hardware paradigms, particularly burst buffers and heterogeneous platform architectures. Data systems performance analysis (to disk and across the interconnect) for both reading and writing data should be one of the driving factors toward whether post-hoc processing is acceptable or *in situ* is required. There is often a threshold before a code team will consider implementing *in situ* analysis. For some cases it was when the simulation spent more than half its time writing to disk. However, there was no consensus on this question.

From a support perspective, the goal is for visualization and analysis use cases to drive the development of necessary *in situ* capabilities, with the visualization and analysis easily specified by domain scientists such that visualization scientists or support personnel are not the primary users. Since *in situ* visualization and analysis is integrated with the simulation code, the *in situ* code needs to be a first-class citizen with the same level of support (both developer and end user) as the simulation code.

In situ analysis does enable some new simulation features, particularly computation steering and experimental acquisition strategies based on forward modeling of instruments. Overall, since hardware and software design is in flux and not set in stone, a whole-system (hardware, software, and support) design is desired.

What visualization facilities and capabilities do you support for both local and remote users?

The primary visualization post-processing software packages supported are VisIt, ParaView, EnSight, and VMD. The primary frameworks for *in situ*/in transit development are Catalyst (ParaView) and Libsim (VisIt) for visualization and analysis, and ADIOS and GLEAN for in transit data movement. The latter could potentially be extended to provide native visualization and analysis, or they could be coupled with an *in situ* library.

The typical visualization/analysis paradigm supported for remote access was typically post-hoc via VNC, NX, and RGS remote displays. There was discussion that *in situ* or in transit processing might make this easier since the data would be reduced via *in situ* analysis to a coarser mesh, data subset, or image.

For remote access, one of the biggest challenges was accessing resources behind firewalls and opening the necessary ports to enable remote users.

What are your major strategies and initiatives over the next 5 and 10 years? How do they affect staffing levels?

There is a common strategy across all of the labs to pursue *in situ* and in transit visualization, and the labs work together on larger initiatives funded by ASC and ASCR on these initiatives. Some individual strategies and initiatives varied across the labs:

- ALCF: “Big Data” Initiative
- LLNL ASC: Data Movement and Data Reduction
- LANL ASC: Integrating Statistics Experts into Visualization and Analysis, Burst Buffers
- SNL ASC: Extreme Scale and Uncertainty Quantification
- NERSC: Application Readiness, Burst Buffers, Data Strategy, and Facilities

What are your current efforts and/or site configuration in this area?

All of the labs have major initiatives in this area. Lawrence Livermore is working on integrating Libsim into its production environment; the multi-institution SDAV (SciDAC Institute of Scalable Data management, Analysis, and Visualization) efforts included *in situ* feature detection and other research efforts. Sandia was integrating *in situ* capabilities into the Sierra suite of simulation codes. Similarly, Los Alamos is integrating *in situ* capabilities into its codes.

What are your mandates and constraints?

Simulation, analysis, and visualization all impose different requirements on the HPC architecture. In an ideal setting the visualization and analysis nodes would be a separate class of nodes from the compute nodes, incorporating differing architectures and perhaps additional resources such as extra memory or burst buffers; however, this is often limited by funding and other pressures. Regardless of node type, it is imperative that the visualization nodes have a high availability for interactive use, especially if data is being post-processed.

How do you forecast future needs and requirements?

The primary methods of forecasting future needs and requirements are other DOE workshops, particularly those centered on domain science needs. In particular, members were interested in how much data and what data processing requirements were required. These workshops call upon visualization experts and especially domain scientists to establish workable requirements.

What opportunities exist for productive collaborations among DOE HPC centers?

Current hardware collaborations include the Trinity/NERSC-8 procurement between ACES (LANL ASC and SNL ASC) and NERSC and the CORAL procurement among OLCF, ALCF, and LLNL ASC. Software collaboration efforts include SDAV in general and the VTK-m collaboration between the VisIt and ParaView teams in particular.

Collaboration gaps were as follows:

- Collaboration tended to be a bottom-up approach. This has its advantages in that the efforts are close to the actual local needs; on the other hand, a balanced approach with coordination between labs allows leveraging common efforts.
- A set of visualization and analysis benchmarks is desired.
- There was little consensus on what we could leverage from industry or the other labs.

What are the biggest challenges and gaps between what you can do today and what will be required in 5-10 years?

Historically, visualization and analysis applications have not been designed to operate within the same set of resource bounds as a simulation code. Often the resource limits were not considered until they became an issue. With in-transit, and especially *in situ*, where the visualization and analysis code is potentially coupled directly into the same executable, the visualization code needs to be mindful of resource limitations. In addition, while these applications are heavily tested, with a tightly coupled code, the visualization and analysis has the potential to bring down the simulation.

As visualization and analysis develops *in situ* and in-transit approaches, currently existing limits on resources will need to be reconsidered. For most *in situ* tasks, a full visualization application is not necessary and will result in a large memory overhead. Work needs to be done to reduce the memory footprint both for the visualization and analysis library and for any data copies necessary to generate the visualization. Furthermore, there needs to be a larger effort toward testing the full functionality exposed for *in situ* visualization and analysis to limit crashes and ensure the validity of results. There is a cost in performance and memory for using an *in situ* or in-transit approach, and those performance and resource tradeoffs between the simulation and visualization/analysis need to be explored, negotiated, and documented. User education and training related to the deployment of *in situ* capabilities is going to be important. In particular, informing users about the tradeoffs between conventional post-processing and *in situ* analysis, and deciding on and implementing canned visualization and analysis *in situ* capabilities, is going to be an iterative process.

Software engineering and co-design will be major issues as *in situ* capabilities are integrated into simulation codes. There have been different levels of coupling in different codes, and the best avenue to take is not clear. The integration effort must be coordinated between the simulation code and the *in situ* code, with the end user driving the use cases. Often the integration falls to whichever group has the funding to perform the integration, but this cannot be performed well without coordination among the simulation code, *in situ* code, and domain scientists.

ADIOS, GLEAN, and potentially HDF5 provide an interesting avenue for *in situ* and in-transit integration. I/O is viewed as an external library that is plugged into the simulation. An in-transit approach could leverage these libraries to plug visualization/analysis into a simulation code.

The user interface for *in situ* and in-transit visualization and analysis will be significantly different from traditional post-processing applications. Current efforts use everything from a standard post-processing GUI to python scripts to low-level integration with the simulation input deck. Each has its benefits and issues and tends to be tailored to different types of users. Python and a post-processing GUI allow for the full range of visualization and analysis operations but tend to be more fragile and harder to test. Input deck integration is easier for domain scientists to understand and easier to fully test, but it does not easily expose the full range of operations.

In situ and in-transit visualization and analysis are also being integrated into current workflow software, including the Sandia Analysis Workbench, UQ applications, and the ACME and CASCADE projects. It is important to maintain the data provenance and ensure the replicability of results regardless of where the simulation and visualization/analysis are run (different node architectures, local vs. remote, etc.).

Exascale architectures will present unique challenges for *in situ* visualization and analysis. Chief among them is the heterogeneous nature of the proposed architectures. This represents a major shift from the typical MPI-everywhere model used for current post-processing and *in situ* visualization and analysis tools. With the increase in concurrency and lower memory per core, the tools and libraries will need to evolve to take advantage of new computing paradigms.

Furthermore, as visualization/analysis moves toward an *in situ* paradigm, contention and resource-sharing issues between the simulation and visualization/analysis will need to be addressed. In particular, the visualization and analysis code needs to be robust, as failure has a much larger impact on a tightly coupled code than in a post-processing paradigm.

Computing centers will need to plan to support *in situ* and in-transit methods as they become more prominent. These technologies are more complicated to use than the traditional post-hoc methods, and centers would be well advised to anticipate this increased labor effort as well as demand for these technologies, which are one primary avenue of dealing with the widening gap between computational and I/O capacity.

Describe some practices that you think are effective, as well as lessons learned that would be helpful to other centers?

Involving domain scientists at the beginning of a research or development effort was seen unanimously as critical to ensuring both that needs are met and that the tools developed are useful and used. The end tools should be designed for the intended audience and not for ease of development or use by visualization specialists.

Often the first implementation of a new technology exposes issues—for example, in more than one *in situ* implementation the application footprint doubled when first integrated. This should not be

cause for immediately discarding the technology, but should lead to further efforts to optimize the implementation. Those efforts are currently under way.

As the integration between visualization and analysis tools and simulation codes tightens, greater effort must be placed on application stability and the integrity of results. For one integration effort, a serious emphasis was placed on automated testing of all aspects of the *in situ* capability. This testing identified several issues and helped to avoid regression errors during development and integration.

Systems implications of *in situ* and in-transit analysis

Depending on the node and machine architecture, in-transit visualization and analysis could require one type of node for simulation and another for visualization and analysis. This leads to issues of co-scheduling (if the nodes are on different machines) or job reservations containing multiple node types (if the nodes are on the same machine).

The addition of *in situ* and in-transit visualization and analysis methods will need to be communicated to system administrators and end users, along with the trade-offs of those approaches compared to traditional post processing. Since visualization and analysis algorithms tend to have different communication paradigms than simulation codes, *in situ* processing may expose errors and issues at the system level that were not present when only the simulation code was executed.

Conclusions

- *In situ* and in-transit analysis and visualization will be necessary components of modeling and simulation at exascale.
- Visualization and analysis capabilities integrated into the simulation code should be driven by the end user and result in a collaboration among the relevant parties (hardware, simulation, visualization and analysis, and the domain scientist) to ensure success.
- While *in situ* and in-transit visualization and analysis are very useful tools, they will not completely replace traditional post-processing capabilities, and these capabilities should also be discussed in the exascale context.

Data Management Policies

Session D1SC

Authors

William Allcock (ALCF), Julia White (OLCF)

Contributors

Dee Magnoni (LANL ASC), Kyle Lamb (LANL ASC), Mark Gary (LLNL ASC), Sasha Ames (LLNL ASC), Jeff Broughton (NERSC), Ashley Barker (OLCF), Bill Collins (SNL ASC), Joel Stevenson (SNL ASC).

Summary

When asked to address the current state of data management policies, representatives of this HPCOR committee concluded that at present, data management at the DOE facilities is largely driven by capacity and I/O performance planning but that other aspects—such as lifetime, provenance, curation, and access—are becoming increasingly important. Policies about data management—from the federal level to the individual facility—are in the nascent stages, and decisions on what capabilities to provide users are being explored. The DOE facilities are taking an active role in helping to identify and shape policies and guidance to enable a data management infrastructure. Ultimately, data will be on an equal footing with computation simulations.

Key Points

- Capacity and I/O performance planning largely drive data management at DOE facilities, but other aspects—such as lifetime, provenance, curation, and access—are becoming increasingly important.
- Policies about data management—from the federal level to the individual facility—are in the nascent stages, and decisions on what capabilities to provide users are being explored.
- Policies about what data is kept, how long it is kept, and how it is accessed will need to adapt. Data must be managed potentially for many years after the project that created it has ended.
- Data access for widespread scientific collaborations will become more important, and this will almost certainly challenge long-standing DOE security policies.
- User teams requesting federal funding are asked to provide data management plans, but few (if any) guidelines have been provided to the facilities regarding resource delivery. Facility services and policies may need to be revised dynamically in response to a rapidly evolving federal consensus.
- Data variety and veracity factors and the enabling factors such as tracking provenance, curation, etc., will drive the most striking changes, including policy, in the DOE HPC data management arena.

Opportunities for Collaboration Among DOE HPC Facilities

- Collaboration is needed at all levels, from individual investigators to principal funding agencies, and the facilities should be major players and drivers in these collaborations: between the facility and federal program managers, between the facilities themselves, between the facilities and major data producers, between facilities and software developers, and between the facilities and the science teams.

Overview

DOE HPC facilities have been, and will continue to be, at the forefront of managing large volumes of scientific and technical data. Historically, the large bodies of data at HPC user facilities are the result of single or multiple simulations produced by research teams awarded compute time. Until a few years ago, only the project members accessed the data, often pulling the file sets to their home institution for analysis, visualization, and archive. This workflow model drove DOE HPC architectural balance, hardware selection, software decisions, and data management policy for the past few decades. We are, however, on the verge of a paradigm shift. Science use cases are demanding and are producing a wider variety of data, with different access patterns and often with community-specific services now part of the workflow. The value (and cost) of data relative to computation is growing and, with it, a recognition that concerns such as reproducibility, provenance, curation, unique referencing, and future availability (when the hardware, software, and data formats may no longer exist) are going to become the rule rather than exception in our scientific communities. Addressing these concerns will impact every facet of facility operations and management if we are to continue to effectively support this increasingly data-focused environment for scientific discovery. Optimal balance of hardware architectures may change. As much or greater emphasis will be given to making software design changes to optimize data movement over changes to improve computational efficiency. Policies about what data is kept, how long it is kept, and how it is accessed will need to adapt. Data access for widespread scientific collaborations will become more important, and this will almost certainly challenge long-standing DOE security policies, which currently stipulate that every user accessing DOE systems be documented and vetted by the facilities.

This nascent shift in user requirements toward broader access—where diverse data is available to investigators beyond the originating project team and for a longer period of time—is reflected in discussions at the federal level. User teams requesting federal funding are asked to provide data management plans, but few if any guidelines have been provided to the facilities regarding resource delivery. At this initial stage, facility services and policies may need to be revised dynamically in response to a rapidly evolving federal consensus. We must walk the line between over-specifying policy too soon and stifling progress and under-specifying policy and having a patchwork solution that is suboptimal and impossible for users to negotiate. The facilities can self-initiate actions to promote and facilitate change, such as educating the science communities about what is changing, why it is changing, and what impact it might have on their community. We may also build system tools, for example, automatically capturing information about details of how software was built on facility systems. But the most important factor in determining the path of change will be collaboration at all levels, from individual investigator to principal funding agencies, and the facilities should be major players and drivers in these collaborations: between the facility and federal

program managers, between the facilities themselves, between the facilities and major data producers (such as the light sources), between facilities and software developers of workflow and data management tools, and last, but certainly not least, between the facilities and the science teams.

The change to a data-focused infrastructure will have a significant impact on staffing and funding for the facilities. One of the significant differences between computation and data is that computation is ephemeral. Researchers run on the machine for a certain amount of time, and once done, that project has minimal further computational demands on the facility. For data to continue to benefit the project and the broader community beyond this point, however, it must be managed, potentially for many years after the project that created it has ended. Traditionally this sharing of scientific discovery is done through papers published in scholarly journals. Taking this a step further and sharing petabytes of data requires specialized resources for archiving and custom programs to access and interpret the data. Delivering these solutions will require staff with specialized expertise. If budgets do not increase substantially, this means stretching and reprioritizing the efforts of existing personnel. The problem is compounded because their expertise is highly valued in the commercial world, driving competition from an industry that can provide substantially greater compensation to address the ubiquitous challenge of “big data.”

While there is no one accepted definition of “big data”, one that is commonly used is known as the “The 3Vs”: volume (the amount of data), velocity (the bandwidth or rate in bytes per second at which the data is accessed), and variety (the types of data and how they are being used). There is also now often a fourth V (veracity}, for how trustworthy or reproducible the data is. DOE HPC facilities have significant experience with the volume and velocity factors. But variety and veracity factors and the enabling factors such as tracking provenance, curation, etc., will drive the most striking changes, including policy, in the DOE HPC data management arena.

Continued discussion among the representatives of DOE HPC facilities is needed to successfully anticipate user needs for data storage, analysis tools, etc. To facilitate these discussions, it was suggested that a data working group could be charged with outlining current and future requirements. In an environment where staffing levels are likely to remain static, the facilities will leverage the “lots of data” expertise with “big data” challenges.

Addendum: Overview of Discussion Points

How is data management planned?

Data management at the facilities is largely limited to the traditional storage characteristics: space, bandwidth, and retention as well as WAN bandwidth. This information is gathered through a combination of sections in proposals and direct discussions with the users. None of the centers have a significant position or presence in the larger data management issues such as provenance, metadata describing the data, or hosting public data sets.

Should centers be part of scientists' data management plans? If so, how?

While all the centers agreed that they should be part of the scientists' data management plans, there was a wide disparity in views on how they should be involved. At present participation is primarily with the traditional storage characteristics (space, bandwidth, retention). For the future, views varied widely. Some of the possible areas of involvement for the facilities included:

- Educating users about what things to consider in a good management plan
- Policies to reward users with good management plans
- Leveraging the expertise and economy of scale of the facilities for supporting community sharing infrastructure. The facilities need to be paid for this, but it should provide a win in terms of cost/performance compared to the communities doing it themselves.

What facilities and policies are in place for data retention and access?

Most facilities will automatically delete (purge) data off disk if it has not been accessed for a specific amount of time, varying from 2-8 weeks. To avoid losing data, users have to move or copy the data to another storage location. Some facilities have a "project" or "campaign" disk storage system that is slower, but larger, where data can be stored for a longer period of time. In terms of tape, all of the facilities were similar in that, so far, tape retention time is infinite, although data growth and costs may force a change to this policy in the future.

What standards for data repositories and archives are in place, and which ones do you plan to support?

All of the facilities support the typical UNIX file system layout: A home file system for things like source code, locally built libraries, etc. and a fast parallel file system for output from codes running on the supercomputers. However, none of the facilities support or have really given any consideration to supporting the repository or archive standards that come with public data repositories such as ISO 14721–Open archival information system (OAIS).

How is access to the broader community provided?

Most of the centers either have, or are moving toward, using the Science DMZ concept, with GridFTP and the Globus Online web service being the primary mover. The classified facilities have other systems for moving data “over the fence” into and out of the classified systems.

How do you balance storage costs with data retention and access policies?

To date, this is largely accomplished by balancing the size and capacity of various storage systems. Small, fast file systems are used for initial data writes, while slower, larger disk systems are used for intermediate-term data storage, and tape is used for long-term archival storage.

What are the challenges and possible solutions?

While capacity is an issue, the real issue is the bandwidth-to-capacity ratio—achieving the right balance. The burst buffer is the current attempt at addressing this issue. The other axis for addressing this issue is reducing the amount of data written to disk. This is being accomplished somewhat through education of users (do you really need to store every attribute at every time step?) but also via the efforts around *in situ* or in-transit processing, which does at least initial data reduction before it is written to disk.

Supporting Data-Producing Facilities and Instruments

Session D1SD

Authors

Stuart Campbell (SNS/ORNL), David Skinner (NERSC)

Contributors

Stuart Campbell (ORNL), Jeff Cunningham (LBNL), Jack Deslippe (NERSC), Rudy Garcia (SNL), John Harney (OLCF), Kaki Kelly (LANL), Galen Shipman (OLCF), David Smith (LLNL), David Skinner (NERSC), Ilana Stern (NCAR), Craig Ulmer (SNL), Dino Pavlakos (SNL), Yao Zhang (ANL)

Executive Summary

There is an increased awareness of the need for BES and BER data-producing experimental facilities to be more tightly integrated to the ASCR HPC facilities. A closer collaboration could decrease time to publication, avoid costly rework, and allow for improved real-time analysis of results. Most importantly, it may also enable new science from the national user facilities.

Several challenges and constraints need to be overcome to maximize the investment in both the DOE HPC and the experimental facilities. First, there is an urgent need for the availability of computational resources to be co-scheduled with the experimental beam time. Second, there is a need for defining standards across facilities and centers, requiring synchronization of data formats, workflows, and environments to allow software for one facility to work everywhere. One such example is to employ a federated authentication system, as often users will make measurements at more than one experimental facility and want to co-analyze the results. Third, there also needs to be expertise (liaisons) at both the HPC facility and the scientific facility to bridge the gap between the two. Staffing is an area of major concern, especially when the labs are competing for talent with Google, Facebook, etc.

It is clear that the strategies for a number of facilities are still in a state of flux for new areas of computing for data-intensive workloads. A number of existing inter-facility partnerships already exist (e.g. NERSC-ALS/LCLS and OLCF-SNS), but these should be strengthened and expanded. We need to build on existing inter-facility partnerships toward generalized services for data-intensive science. The HPC centers will need to have computational stewardship of both traditional HPC and the emerging portfolio of computing, data, and network needs. It is essential to draw expertise from both the computational/data-intensive (ASCR) side and the data-producing scientific facility (BES/BER) side. This will require investment from different areas within DOE, such as ASCR, BES, and BER.

Key Points

- A tighter integration between BES and BER experimental facilities and ASCR HPC centers would decrease time to publication and enable new science from the national user facilities.
- There is a growing need for new data services not traditionally provided by HPC facilities, such as long-term stewardship of valuable experimental and observation datasets, data analysis systems optimized for mining, and visualization of these data.
- HPC centers will need to have stewardship of both traditional HPC and the full portfolio of computing, data, and network needs.
- It is essential to draw expertise from both the computational/data intensive (ASCR) side and the data producing scientific facility (BES/BER) side.
- It is unknown whether or not upcoming exascale platforms will be able to satisfy data-intensive computing needs or if dedicated platforms will be required.
- Computational facilities need to be fully engaged with the DOE experimental facilities' current and future plans for their data and compute needs.

Opportunities for Collaboration Among DOE HPC Facilities

- DOE should build on existing inter-facility partnerships to produce generalized services for data-intensive science.
- There is a need for standards across facilities and centers to enable synchronization of data formats, workflows, and environments in a secure and transparent manner.

Discussion

Classification of Data-Producing Facilities

Data-producing facilities can be classified into one of three types:

1. “Bursty” facilities that have a very high data rate over a short period of time, for example the Linac Coherent Light Source at SLAC, some individual (neutron/x-ray) scattering beam lines, telescopes, and Community Earth System Model generation
2. “Constant stream of data” facilities that, over a given period of time, have a data rate that is approximately a constant (for example, SNS, ALS, APS, JGI, ARM)
3. “Distributed with aggregation” instruments that collect data over a period then are collected together (for example, sensor networks, fluxnet, CDIAC)

Of course it is important to note that the differentiation between the first two categories will largely depend on the time scale under consideration. It is entirely possible to have a “bursty” instrument at a constant data rate facility that would give a spike of data on top of a flat background.

What are your major strategies and initiatives over the next 5 & 10 years? How do they affect staffing levels?

It is clear that the strategies for a number of facilities are still in a state of flux for new areas of computing for data-intensive workloads. We need to build on existing inter-facility partnerships toward generalized services for data-intensive science. The HPC centers will need to have computational stewardship of both traditional HPC and the emerging portfolio of computing, data, and network needs. Staffing is an area of major concern. It is essential to draw expertise from both the computational (ASCR) side and the data-producing (BES/BER) side. This will require and investment from different areas within DOE, such as ASCR, BES and BER.

What are your current efforts and/or site configuration in this area?

OLCF

- ADARA² – Streaming data directly from SNS to HPC resources
- CAMM³ – Integration of simulation with experiment
- CADES⁴ – Cross-cutting effort to provide lab-wide data services and domain expertise
- ACME⁵ – End-to-end testbed for production runs of the DOE Climate and Environmental Sciences Division of BER
- CDiac⁶ and ARM⁷ Archives – long-term archival of observational climate data
- ASCR/BES Data Pilot Project⁸ (APS/SNS)
- IFIM⁹ – workflow, analysis, and storage for this initiative
- ALICE¹⁰
- Panda¹¹ workflow on HPC

LLNL ASC

- ALICE¹⁰
- PCMDI¹² climate archive

LBNL

- LCLS-II pilot work¹³ (Photon Science Speedway)
- ASCR/BES pilot⁸

² Accelerating Data Acquisition, Reduction, and Analysis (<http://techint.nccs.gov/data.html>)

³ Center for Accelerating Materials Modeling (<http://camm.ornl.gov/>)

⁴ Compute & Data Environment for Science (<https://ornlwiki.atlassian.net/wiki/display/CADES/>)

⁵ Accelerated Climate Modeling for Energy (<http://climatemodeling.science.energy.gov/>)

⁶ Carbon Dioxide Information Analysis Center (<http://cdiac.ornl.gov/>)

⁷ Atmospheric Radiation Measurement (<http://www.arm.gov/>)

⁸ http://science.energy.gov/~media/ascr/pdf/research/scidac/ASCR_BES_Data_Report.pdf

⁹ Institute for Functional Imaging of Materials (<http://www.ornl.gov/science-discovery/advanced-materials/research-areas/institute-for-functional-imaging-of-materials>)

¹⁰ <http://aliceinfo.cern.ch/>

¹¹ <https://twiki.cern.ch/twiki/bin/view/PanDA/PanDA>

¹² Program for Climate Model Diagnosis and Intercomparison (<http://www.pcmdi.llnl.gov/>)

¹³ <http://cs.lbl.gov/news-media/news/2014/photon-speedway-puts-big-data-in-the-fast-lane/>

- CXIDB¹⁴
- NCEM¹⁵
- SPOT Suite¹⁶ (ALS/SSRL/LCLS Simulation and Analysis Framework, Portal)
- Daya Bay Neutrino Experiment¹⁷
- 20th Century Re-Analysis Project¹⁸
- Palomar Transient Factory¹⁹
- NEWT Web API²⁰, FireWorks²¹ workflow manager
- PDSF²², JGI²³

ANL

- Integrating Simulation and Observation: Discovery Engines for Big Data²⁴
- ASCR/BES pilot⁸ (SNS/APS)

NCAR

- GLADE²⁵ – PB scale disk
- ESG²⁶ – publication of Community Earth System Model (CESM) runs

What are your mandates and constraints?

The OSTP memos and Office of Science directions in Open Access provide a clear indication that the management of data needs to be planned. This mandate impacts facilities and the PIs who apply for resources at the facilities.

There is a growing expectation from users that any solutions developed by the facilities are based upon (or developed as) open source (OSS). It is important that the facilities can provide the connection between the HPC and data resources and the experimental facilities while still maintaining their security policies. There is a growing need to provide queuing for both traditional batch and near real-time computing, for example when compute resource allocation needs to be co-allocated with experimental beam time. There is a growing need for new data services not traditionally provided by HPC facilities, such as long-term stewardship of valuable experimental and observation datasets, data analysis systems optimized for mining and visualization of these data, and experts in the use of these technologies to partner with experimental and observation facilities to develop solutions.

¹⁴ Coherent X-ray Imaging Data Bank (<http://www.cxidb.org/>)

¹⁵ National Center for Electron Microscopy (<http://foundry.lbl.gov/facilities/ncem/>)

¹⁶ <http://spot.nsls.gov/>

¹⁷ <http://www.bnl.gov/science/dayabay.php>

¹⁸ http://www.esrl.noaa.gov/psd/data/20thC_Rean/

¹⁹ <http://www.ptf.caltech.edu/>

²⁰ <https://newt.nsls.gov/>

²¹ <http://pythonhosted.org/FireWorks/>

²² <http://www.nsls.gov/users/computational-systems/pdsf/>

²³ DOE Joint Genome Institute (<http://jgi.doe.gov/>)

²⁴ <http://www.mcs.anl.gov/project/integrating-simulation-and-observation-discovery-engines-big-data>

²⁵ Globally Accessible Data Environment (<https://www2.cisl.ucar.edu/resources/glade>)

²⁶ <https://www.earthsystemgrid.org/home.htm>

How to do you forecast future needs and requirements?

Computational facilities need to be fully engaged with the DOE experimental facilities' current and future plans for their data and compute needs. This will mean being engaged with major facility and instrument upgrades during the planning phase to understand computing and data requirements and how best to structure new services to meet the needs of these facilities.

What are the biggest challenges and gaps between what you can do today and what will be required in 5-10 years?

There is some uncertainty around the computer architecture for emerging data-intensive spaces. One question that remains unanswered is whether the upcoming exascale platforms will provide all that is needed to support data-intensive computing, or will we need dedicated platforms?

What opportunities exist for productive collaborations among DOE HPC centers?

It is clear that some interesting opportunities for productive collaborations exist between DOE HPC facilities. Most facilities operate in a stand-alone mode, so the opportunities that have been uncovered so far are those in which a bottleneck or data-scaling issue has necessitated fast networks and big computers to address the data deluge. Using these pilots to motivate a broader framework for multi-disciplinary science that combines experiment, simulation, and theory is a longer term goal.

There is a real need to define standards across facilities and centers in terms of data formats, workflows, and environments to allow software written for one facility to work everywhere. The users of the various facilities often get frustrated at the number of authentication hops that are required, so an opportunity exists to employ some sort of federated authentication system. Users will often make measurements at more than one experimental facility and want to be able to co-analyze the results. Some steps are being made in this area (for example, the ASCR/BES Pilot Project between SNS/OLCF and APS/ANL), but there is room for more collaboration.

Infrastructure

Session D2SA

Authors

Christopher Beggio (SNL ASC), Robin Goldstone (LLNL ASC)

Contributors

Bill Allcock (ALCF), Christopher Beggio (SNL ASC), Stuart Campbell (ORNL), Clay England (OLCF), Doug Fuller (OLCF), Robin Goldstone (LLNL ASC), Jason Hick (NERSC), Kyle Lamb (LANL ASC), Jim Silva (LLNL ASC), Jay Srinivasan, (NERSC), Ilana Stern (NCAR), Craig Ulmer (SNL ASC), Venkat Vishwanath (ALCF)

Abstract

The HPCOR review addressed major areas of technical and administrative concern to data-driven science applications, including system configuration, visualization, data management policies, data production and instrument support, infrastructure, user training, workflows, and data transfer. This report summarizes proceedings and findings of the Infrastructure breakout session.

Introduction

The HPCOR review convened June 17-19, 2014 in Oakland, CA, in conjunction with the Joint Facilities User Forum on Data-Intensive Computing. The purpose of this conference was to assemble members of the DOE Office of Science and National Nuclear Security Administration laboratories that are responsible for the planning, architecture, and operation of computing platforms to support data-driven science. HPCOR participants were organized into collaborative teams in breakout sessions to address major areas of technical and administrative concern, including system configuration, visualization, data management policies, data production and instrument support, infrastructure, user training, workflows, and data transfer.

Key Points

- Facilities have the ability and infrastructure to ingest and write data to storage, but more work and capability is needed to perform efficient random reads and queries for data-intensive operations.
- Facilities also have requirements for an improved software middleware layer beyond hierarchical storage management. This middleware must be locality aware and include integration with job scheduling.
- Driven by a void in formal budgeting for big-data type systems, many sites leverage existing HPC installations when possible. But since existing HPC systems are not universally architected with data-intensive applications in mind, separate or additional data-intensive

processing segments must be designed, purchased, and built to accommodate node-local storage, burst buffer, and unique software stack requirements.

- The continual tug-of-war between security and data movement at the network edge challenges the improvement of data services.
- Some sites have deployed private cloud architectures effectively. However, public cloud offerings are not tailored toward “largest-scale” data-intensive and data analytics processing. Their use creates availability, reliability, performance, and security concerns for the national laboratory complex.

Opportunities for Collaboration Among DOE HPC Centers

- There are opportunities for continued collaboration among sites that force multiple talent and resources into generalizable solutions and contributions to industry. There are also opportunities to leverage mature solutions from advanced sites to minimize duplicative efforts and expenditures.
- Personnel across sites are actively investigating available options and alternatives to existing storage technologies and methods.

Discussion

During the Infrastructure (D2SA) session, contributors responded to a number of questions related to the identification, design, operation, and maintenance of infrastructure components that support data-driven computing platforms. The following questions provided the framework for the specific questions and discussion among contributors:

- What supporting infrastructure is needed to enable data-driven science?
- Which of these play primary roles in supporting your center's data-driven science: networking between resources, shared or local disk, archival storage, science data gateways or portals, consulting, and databases?

Questions and Responses

How does data-driven science impact your HPC network and storage architectures (including archive)? Are significant changes required, or do you just need "bigger" pipes and "more" storage?

Due to extensive experience with traditional HPC, facilities have the capability and infrastructure to ingest and write data to storage, but more work and capability is needed to perform efficient random reads and queries for data-intensive operations. Facilities also have requirements for an improved software middleware layer beyond hierarchical storage management. This middleware must be locality aware and include integration with job scheduling. Many of the solutions to data-intensive

computing applications are site-specific and do not always lend well to canned or standardized solutions. Despite this, sites would prefer to leverage industry standards and generalized solutions whenever feasible, to adopt best practices and minimize time to production.

How is data-driven science impacting your HPC planning processes? Have you formally integrated data science requirements into your planning process or are these efforts currently being pursued "on the side?"

Driven by a void in formal budgeting for big-data type systems, many sites are accomplishing deployments of these types of systems through the reapplication of hardware. The increasing need for dedicated data-intensive computing platforms with integration into HPC systems is being recognized and addressed through formal programs and procurements. There is a preference among sites to leverage existing HPC installations where possible by implementing visualization and data-intensive processing segments in lieu of separate, monolithic platforms. Since existing HPC systems are not universally architected with data-intensive applications in mind, separate or additional data-intensive processing segments must be designed, purchased, and built to accommodate node-local storage, burst buffer, and unique software stack requirements.

How is data-driven science affecting your HPC procurements? Is the spending balance shifting more toward storage and alternative data-centric architectures? Is more money needed/available to meet these evolving requirements?

Due to the differences in mission and requirements across Office of Science and NNSA facilities, there is little coordination and uniformity of data-intensive platforms. Visualization platforms have been re-tasked to perform data-intensive and data analytics type applications due to larger memory footprints, I/O performance, node-local storage, and display capabilities. Evolution of use-case definition and a better understanding of data-analytics requirements will promote the architecture and design of purpose built systems, as contrasted with repurposed systems. The need for and benefit of burst-buffer technology in data-intensive computing applications is recognized across many sites. However, this will be at the expense of size and quantity of other HPC system components unless additional funding is allocated.

What alternative storage technologies and file systems are being considered to support big data storage requirements?

Personnel across sites are actively investigating available options and alternatives to existing storage technologies and methods. Among these alternatives are burst-buffer solutions, including node-local solid-state storage devices, PCI-e attached storage, and memory bus non-volatile memory technologies. Solid-state storage is also being considered in other parts of data processing clusters, such as I/O nodes and storage and gateway nodes. Branded vendor products are also being considered to solve the unique memory and storage requirements, such as EMC Viper, HP SL-series, and IBM Ultra DIMM. Additionally, many sites are exploring how to improve storage performance

by replacing enterprise-class disk storage systems based on high-overhead RAID algorithms with JBOD disk arrays coupled with intelligent block replication and checksum verification software. File system alternatives are also being investigated at many sites, including ZFS, Lustre on ZFS, storage appliances with embedded file systems, and processing such as Xyratex, Ceph, and Gluster.

How are WAN providers responding to inter-facility large-scale data movement requirements?

Site personnel recognize the need for network capacity to increase commensurate with data quantity, an increase in slope of data creation capability, and further future expansion of both. The continual tug-of-war between security and data movement at the network edge challenges the improvement of these services. The combination of network packet scanning and encryption for virtual private networks frequently places purpose-build devices at the point-of-presence that limit the ceiling for network transmission, regardless of total connection bandwidth. Furthermore, planning for wide-area network connectivity and expansion is often performed independently of high-performance and data-intensive computing platform planning, creating the potential for function and performance disconnection.

Do data-centric workflows impact your security posture—for example, need for firewall bypass to efficiently transfer large data flows across security boundaries, real-time data acquisition, etc.?

The science DMZ model has been implemented and is a working solution for many sites. Due to user access requirements across security domains, there is a need for federated authentication and authorization solutions. Due to associated risks, many sites are dissuaded or prohibited from exposing production data processing resources to the Internet. Therefore, some may be forced to isolate resources in an established DMZ, creating problems with federated user access and storage duplication.

Are commercial cloud and big data solutions applicable to the scientific community, or do we need to "roll our own?"

Some sites have deployed private cloud architectures effectively. Public cloud offerings continue to have data integrity and security concerns as well as reliability issues that are generally not acceptable for secure sites and sensitive data. Many enterprise service offerings are benefitting from public cloud architectures such as remote collaboration, open source software projects, and publicly released literature publishing. Commercial solutions generally do not target large-scale scientific computing and fall short in capacity and performance requirements. The lack of public cloud offerings tailored toward "largest-scale" data-intensive and data analytics processing creates availability, reliability, performance, and security concerns for the national laboratory complex.

Findings

Opportunities

It is evident from the discussion that all sites are exploring solutions for a similar set of problems. There are opportunities for continued collaboration among sites that force multiple talent and resources into generalizable solutions and contributions to industry. Furthermore, there are also opportunities to leverage mature solutions from advanced sites to minimize duplicative efforts and expenditures.

Best Practices

There is a significant amount of knowledge and documentation surrounding the existing institutional HPC service offerings. This knowledge base should be used whenever practical. The laboratories and user facilities should be prepared to extend existing knowledge of high volume and velocity data storage systems into this newer knowledge domain.

Challenges

The laboratory complex is facing a wide range of requirements related to data acquisition, storage, processing, and movement. The development of talent proficient in data science and analytics results in personnel harvesting by private industry. Presently, the laboratory complex does not receive funding specific to data sciences and analytics. Funding is presently diverted from institutional funding and program-driven HPC.

User Training for Data-Intensive Science

Session D2SB

Authors

Fernanda Foertter (OLCF), Tim Fahey (LLNL)

Contributors

Blaise Barney (LLNL), Karen Haskell (SNL), Bob Balance (SNL), Richard Gerber (NERSC), Dee Magnoni, Glenn Lockwood (SDSC), Ashley Barker (OLCF), Kjiersten Fagnan (NERSC), David Karelitz, Richard Coffee (ANL)

Introduction

There has been an increase in the number of projects and people doing data-intensive computing at HPC centers. Numbers and sizes of datasets, whether from experimental facilities or simulated, are increasing at an unprecedented rate. With this growth comes the need to introduce users to tools and best practices on data-intensive computing. Therefore training will require increased focus beyond using or building generic single-purpose tools, but also a holistic approach to training on the topic of data, from creation to analysis to curation. These challenges are not limited to large HPC centers and naturally create an opportunity for increased training-related coordination between centers of all sizes.

Key Points

- Users currently face a multitude of new challenges surrounding large data sets.
- Many users are unfamiliar with techniques and best practices for managing and analyzing data and need training on techniques, use of tools, and efficient use of hardware.

Opportunities for Collaboration Among DOE Centers

- Many users have access to multiple centers and need to be trained on how to work collaboratively with data across sites.
- DOE centers have many pieces of common data infrastructure, both software and hardware.
- DOE centers can leverage expertise across sites for training on common resources.
- A training collaboration among sites has already formed and is conducting training classes open to users from all centers.

Current Challenges Faced by Users

Users currently face a multitude of challenges surrounding large datasets, for example:

- Bringing data into, and taking data out of, centers
- I/O during simulation and analysis: cache levels, ram disk, burst buffers
- File systems: distributed, silo'ed, lack of FS specific libraries, file striping
- Data archiving: tape drives, data management plans
- Sharing data: with collaborators, with community
- Data analysis: lack of scalable tools, *in situ* vs. post processing

In many ways these issues are interrelated, and the lack of broad guidance means many users continued to perform the workflow before their datasets had grown large. Below are some general areas where collaborative training could improve these issues.

Data transfers in/out of data centers

Despite advances in computing power, networks have not kept pace with data demands of today's scientific computing needs. In some ways, centers have mitigated data transfer issues by building dedicated networks (ESnet, Internet2). However, in today's distributed and collaborative environment, many users have data stored across multiple centers, as well as on local resources that are outside of the dedicated networks. Therefore, they must deal with moving data on networks of varying capability, capacity, and reliability. Users have reported having to monitor individual data transfers, sometimes being prevented from doing bulk transfers due to connection time outs, or use protocols that are far slower.

Opportunities for training include:

- Encourage users to keep data local to where it is computed, and analyze where data is created when possible
- Expose users to tools that automate data transfers, such as Globus
- Inform users of dedicated networks and tuning opportunities

Opportunities for collaboration among centers include:

- Installing dedicated data transfer nodes and training users to use them
- Ensuring data transfer tools work in an optimized way between centers and that users are aware of the proper use of these tools

I/O performance during simulation and analysis

With centers moving to manycore architectures, systems have increased levels of "cache tiers." With large datasets, the bottleneck is often cache misses from system memory to core memory. Future systems are also moving toward larger nodes and the use of SSDs as even bigger caches. This fast/slow memory means explicit and deliberate data movement must be integrated into applications.

Opportunities for training include:

- Data pattern profiling tools
- Transforming data structures to avoid cache misses
- Use of libraries and compiler optimizations

Opportunities for collaboration among centers include:

- Leveraging expertise across centers
- Avoiding duplication of effort by co-hosting or recording training events

File systems

HPC file systems are large and behave differently from the local hard drives most users are familiar with. Often applications will scale poorly if their I/O patterns are not modified before running in a large system.

Opportunities for training:

- Porting language-specific calls to I/O libraries
- Educating users on basics of large-scale file systems, how they behave under load, and their basic topology
- Teaching about file striping and the use of different file formats
- Encouraging the use of workflow tools and practices in places with site-wide file systems

Most centers run similar, if not identical, file systems, and this area is fertile for collaboration between centers. Therefore, coordinating the availability of tuned I/O libraries (such as ADIOS and MPI I/O) and leveraging each center's expertise for the benefit of users when file systems differ should be considered a best practice.

Data archiving and sharing

With the advent of recent data management plan requirements (DMP), users are being required to plan to make available, to their community, data used at DOE facilities after they have generated publications. But even before this happens, projects that last multiple years require archiving of temporary datasets in tape drives since file systems are reserved for active data only.

Opportunities for training:

- Efficient data transfer in/out of tape based file systems
- Efficient use of archival storage systems
- What the center can and cannot provide regarding DMP
- DMP tools for creating plans

DMPs are an evolving topic, and it remains to be seen how publications that rely on the analysis of very large datasets will be required to comply with such mandates and whether centers will need to play a larger role. To that end, training provides an opportunity for requirements gathering within communities that will need help complying with DMP requirements.

Data analysis and visualization

Large datasets have increased dramatically the time users spend on analysis. Unlike production applications that can be scrutinized for their computational readiness, analysis applications have different computational requirements and have lagged behind in efficiency and performance.

Opportunities for training:

- Increased threading and scalability in analysis applications
- Memory and I/O efficiencies
- *In situ* data analysis
- Data analysis tools (pdbR, python-based, domain-specific)
- Visualization tools (ParaView, VisIt)

Opportunities for collaboration among centers include leveraging expertise and sharing knowledge of new tools that become available. An opportunity for center collaboration is the preservation of technical knowledge of tools, their strength and weaknesses, bugs and updates, etc.

User education and training related to the deployment of *in situ* capabilities is going to be important. In particular, informing users about the tradeoffs between conventional post processing and *in situ* analysis, and deciding on and implementing canned visualization and analysis *in situ* capabilities is going to be an iterative process.

Discussion

In the coming decade it is expected that hardware complexity will only increase and with it the need for larger developer teams with more specialized skills among projects using HPC facilities. Therefore training needs to adapt to the inevitable increase in diversity of skillsets among project teams, which will range from purely scientific users of community applications to computer science developers of libraries and tools. Traditionally, training provided to users has been mostly documentation and lecture style. More centers are beginning to offer more hands-on classes, remote interactive webinars, and short courses intended to reach a wider and more diverse audience. This mixed approach has been shown to be popular, but metrics do not exist yet to measure its success. Quantification of the effectiveness of training delivered is a gap that still requires exploration.

Another issues is that, effort is often duplicated at multiple centers that cover similar topics, although this can be avoided with increased collaboration. In addition, reuse of shared didactic materials will lead to step-wise improvement in the information disseminated, preserving the knowledge gained with each delivery. Since experts are spread across centers, it's important to build a database of expertise to ensure the coverage of best practices among topics. Building this network of support staff would also lay the foundation for sharing the effort of training and documentation material between DOE and centers outside DOE.

A monthly call of interested centers to communicate training events is a low hanging fruit that has already developed between DOE labs. In the near future, this cooperative work may require more sophisticated infrastructure to make it easier to share content among centers. This is especially important to draw centers outside the DOE Open Science centers, such as NNSA, NSF, and DOD. This centralized infrastructure could host draft documentation, instructional material, videos, code

examples, information about centers' hardware details, and a communication platform between support staff across centers. However, the complexity of setting up such an infrastructure would require some initial funding to create and some marginal amount of staffing time to ensure availability and accessibility.

Workflows

Session D2SC

Authors

Shreyas Cholia (NERSC), Kevin Harms (ALCF)

Contributors

Yao Zhang (ALCF), Laura Monroe (LANL ASC), Sasha Ames (LLNL ASC), Jeff Long (LLNL ASC), Norbert Podhorszki (OLCF), Andy Wilson (SNL ASC), Rudy Garcia (SNL ASC), Laura Biven (DOE), Stephen Bailey (LBNL)

Key Points

- Workflows—sequences of coordinated compute and data-centric operations—are important for supporting uncertainty quantification (UQ), computational steering, parameter studies, and the near real-time processing of “big data” generated from instruments like light sources, detectors, telescopes, etc.
- Although workflows have been around for a long time, their use on modern HPC architectures is still new. They require documentation, support infrastructure (for example, databases, master control nodes, and virtual environments), and staff to facilitate site integration.
- Security policies, along with authorization and authentication systems, are one of the major stumbling blocks for workflow systems.
- Current HPC resource managers cannot handle the number of jobs required by some of the workflow systems. Many workflow systems were created to overcome this specific issue.
- Existing workflow systems generally require software that hasn’t been part of traditional HPC software stacks. Some HPC centers are exploring allowing users to run virtual machines to accommodate their workflows that may have been developed on another specific platform.
- There has been little standardization on workflow software, partly because no common workflow systems are supported at the different DOE lab sites.

Opportunities for Collaboration Among DOE HPC Centers

- Organize a “Birds of a Feather” (BoF) or workshop as part of a major conference such as Supercomputing.

- Build a catalog of workflow tools along with their capabilities to be able to provide sites and users with guidance in terms of what tools to choose for their jobs.
- Define a set of metrics to evaluate workflows, and offer recommendations based on various criteria such as performance, throughput, ability to handle different classes of problems, feature sets, ease of use, etc.
- Workflow systems will play a key role in connecting resources at different sites. The hope is that something along the lines of a DOE-wide Scientific Data Facility (SDF) will help drive improvements in cross-site security access and tools that will enable multi-site workflow opportunities.

Introduction

This is a summary of the Workflows breakout session at the 2014 HPCOR meeting. The meeting was intended to cover processes and practices for delivering facilities and services that enable high performance data-driven scientific discovery at the DOE national laboratories.

Specifically, this session addressed the role of workflows in data-intensive computing. Participants from several DOE labs and partner institutions were able to offer perspectives on their experiences with workflows and workflow tools across multiple scientific domains. The session was aimed at discussing the following questions:

- What is a workflow?
- What workflow systems are being used?
- What works well and what is missing?
- What infrastructure and support is required?

Defining Workflows and Workflow Systems

Defining a workflow is the first key step in understanding how workflows will be important for future data-driven science projects. Workflows (or more generally, work orchestration) can be simply defined as sequences of coordinated compute and data-centric operations. For example, a workflow can be used to automate a set of applications executing in some defined order to generate data for future steps or process data from previous steps.

Elements of a workflow fall into a handful of categories:

- bag of tasks (Directed Acyclic Graph, or DAG)
- *in situ* analysis
- map-reduce
- provenance tracking
- data movement

Workflow systems have been developed to address the evolving needs of scientific computing as data-intensive jobs become increasingly important. Workflow systems provide a framework to orchestrate and run workflows and can facilitate the execution of various workflow elements. They may also perform data movement within/between systems and automatic provenance capturing. Another aspect common to many workflow systems is the ability to support very large numbers of independent but coordinated tasks (on the order of 1 million). Current HPC scheduling systems do not support efficient queuing and execution of large numbers of jobs.²⁷

Use Cases and Examples

Workloads with high volumes of tasks are common in fields like uncertainty quantification^{28, 29} and the near real-time processing of “big data” generated from instruments like light sources, detectors, telescopes, etc. Workflow systems must be able to support these types of massively data-parallel tasks.

Uncertainty quantification requires running a large number of usually smaller sized jobs to investigate the results of parameter changes. The workflow system allows the automation of submitting those jobs, perturbing the parameter(s) for the needed runs, and tracking and publishing the results.

Near real-time data from large scientific instruments often calls for “small” processing tasks on large amounts of data. These processing tasks can often proceed independently. This type of problem easily translates to the workflow model where many tasks analyze some bit of data and then contribute output information that will be aggregated by later steps in the process. For example, NERSC is using this model to handle beamline data generated from users of the Advanced Light Source, which allows scientists to analyze such data in near real-time.³⁰

We also considered workflows used in simulation runs. As an example, two workflows implemented with the Kepler workflow system are used by the Center for Plasma Edge Simulation at OLCF for monitoring and guiding a particular simulation.³¹ The simulation starts with a particle code that advances the simulations. The workflow system runs a second monitoring code to evaluate the conditions of the simulation; when the flow becomes too turbulent, the particle code is stopped and an MHD code is started. The workflow system continues monitoring the output and when the system becomes less turbulent, the MHD code is stopped and the particle code is resumed from that point. The workflow also runs a system of tools that create and store images from the data generated during the simulation. The provenance data (source, inputs, environment, etc.) relevant to the run is also recorded and bundled with the images created. This data is then moved over to a system that

²⁷ “Cramming Sequoia Full of Jobs for Uncertainty Quantification”, <http://computation.llnl.gov/cramming-sequoia-full-jobs-uncertainty-quantification>

²⁸ J. A. Hittinger, B. I. Cohen, R. I. Klein, “Uncertainty Quantification in the Fusion Simulation Project Verification and Validation Activity”, LLNL-TR-458089, October 5, 2010

²⁹ Tannahill, J.; Brandon, S. T.; Covey, C. C.; Domyancic, D. M.; Garaizar, X.; Johannesson, G.; Klein, R. I.; Lucas, D. D.; Zhang, Y. “The Climate Uncertainty Quantification Project at Lawrence Livermore National Laboratory: I. Initial Analysis of the Sensitivities and Uncertainties in the Community Atmosphere Model”, American Geophysical Union, Fall Meeting 2010, 12/2010

³⁰ “Big Data Hits The Beamline”, <http://cs.lbl.gov/news-media/news/2013/big-data-hits-the-beamline/>

³¹ S. Klasky, B. Ludašcher, and M. Parashar, “The Center for Plasma Edge Simulation Workflow Requirements,” in 22nd International Conference on Data Engineering Workshops (ICDEW’06). Atlanta, GA, USA: IEEE Computer Society, 2006, p. 73.

can be viewed through typical web services.

Major Initiatives and Strategies

The DOE HPC centers see support for workflow systems as critical to the success in the evolving “big data” space. All of the main DOE HPC centers are engaged in support for these workflow systems and trying to find solutions to the problems facing these systems. Each lab has a different particular driver for looking at these systems:

- The ALCF has seen an uptick in the request from users to run a set of smaller jobs running the same application in a larger partition, while a steering code runs independently evaluating job outputs to generate new inputs for the next stage of computation. The ALCF is considering allowing users to run virtual machines on its VM infrastructure so they can run the master process on that.
- The primary driver for LLNL ASC is uncertainty quantification. There are $O(100,000)$ to $O(1,000,000)$ jobs that need to be run and analyzed based on permutations of input parameters.³²
- The OLCF has the CADES initiative, which plans to tie the experimental resources into the LCF for computation support. This initiative will need workflow support for transferring data back and forth as well as starting the various analysis pipelines needed for experimentation results.
- NERSC is being driven by a number of projects that require connecting DOE instruments and facilities with computational resources, and being able to process a high data volume to support these instruments. This includes support for ALS beamlines, telescope data from large astronomical surveys, sequencing data from the Joint Genome Institute,³³ and bioimaging data from NCEM. NERSC currently supports multiple workflow systems and has official documentation for the software. NERSC is actively pursuing more support for these systems to enable users to more easily take advantage of the available compute resources.

Scientists are also creating initiatives around workflows systems. The increasing pressure on storage systems and possible heterogeneous compute systems is leading scientists to develop their own workflows to reduce data movement and the amount of data written to permanent storage.

We were able to identify a useful set of strategies for workflows based on the approach taken by the Scientific Data Group at CSMD (OLCF):

In Situ Analysis: Current practice of postponing analysis/visualization to the post-mortem of a simulation is not viable at exascale because of the prohibitive cost of staging and storing on disk all the data generated by a simulation. Running analysis when and where the data is being generated (ci analysis) becomes a necessity rather than an option.

³² Domyancic, D. and Nimmakayala, R. “Creating a Path to Petascale-Class Uncertainty Quantification Studies,” <http://computation.llnl.gov/system/files/downloads/Computation-Annual-Report-2013-LLNL.pdf>

³³ NERSC JGI JAMO tool <https://www.nersc.gov/news-publications/news/nersc-center-news/2013/new-metadata-organizer-and-archive-streamlines-jgi-data-management/>

Adaptive Analysis: Scientists should be able to interact with and modify the analysis based on analysis results. For instance, if an interesting or unexpected pattern is detected, the scientist should be able to spawn a new diagnostic code and/or instantiate a new feature detection method coupled with a robust classification algorithm to monitor and explore the pattern more closely.

Near Real-Time Analysis: Data analysis should match the rates at which simulation data are generated; on exascale systems, data rates are expected to increase significantly. Scientists should be able to inspect results on the fly and make near real-time decisions to modify the analysis or control the simulation.

Best Practices

Although workflows have been around for a long time, their use on modern HPC architectures is still new. At this point it would be premature to nominate anything as a validated best practice. Here we cover some successful strategies that have been used, as well as future ideas that may be successful but still need to be tested.

Documentation

Given the wide variety of workflow tools and HPC architectures, it is prudent to have workflow systems well documented. Providing a list of available workflow tools and their features is essential to helping users successfully execute work with these systems. This should help mitigate the syndrome in which each user creates a one-off workflow system to support their work.

Support Infrastructure

Many workflow systems require a “master service” or other components to support the workflow system. For example, a workflow system may need access to a database or message queuing service to coordinate its activities. These components often have a larger set of dependencies and requirements that make them difficult to run on individual compute elements. A possible solution to this issue is the use of virtualized environments. These environments allow more flexibility in installing software packages and can optionally be placed on the “border” between internal compute networks and external routable networks.

Site Integration with HPC Staff

The various workflow systems can have rather complex requirements related to software stacks, network connectivity, operating systems, and security policies. These issues necessitate that HPC staff be familiar with workflow systems and actively work with both the developers and the users to integrate the software into the HPC site.

Challenges

Workflow systems provide many challenges and opportunities for the DOE labs regarding adoption and usage.

Security

Security policies, along with authorization and authentication systems, are one of the major stumbling blocks for workflow systems. Workflow systems require automated software to run and interact with scheduling systems, data transfer mechanisms, and software running on the compute hardware. Security policies and systems typically don't allow this type of flexibility. Existing solutions using shared common authentication and authorization services could be adopted to allow more system usage flexibility while still maintaining proper security.

This could be an area for future DOE laboratory collaborations related to adopting shared security frameworks. Federated identity across the DOE complex and the ability to authenticate across domains are key issues that need to be addressed, especially with regard to workflow systems that need to coordinate data from facilities and instruments with HPC resources across multiple sites.

Resource Managers

Current resource managers cannot handle the number of jobs required by some of the workflow systems. Several sites, including the ALCF, LLNL ASC and NERSC, have attempted to submit large quantities of jobs only to find the scheduler either crashed or became too slow in handling the large number of jobs. Many workflow systems have been created to overcome this specific issue.

Non-traditional Software

Existing workflow systems generally require software that hasn't been part of traditional HPC software stacks. User-oriented services needed by workflow tools such as web servers, data publication tools, database servers, and "master" servers are not always available or supported in HPC centers.

Standardization

There has been little standardization on workflow software. There are many software packages that provide similar capabilities. There are also many one-off solutions that have been written for a particular software package and/or compute system. There may be many reasons for this, but chief among them is that no common workflow systems are supported at the different DOE Lab sites.

Next Steps and Opportunities

The logical next steps to move workflow systems forward on DOE HPC systems is to organize a

“Birds of a Feather” (BoF) or workshop as part of a major conference such as Supercomputing.

We also propose building a catalog of workflow tools and their capabilities to provide sites and users with guidance in terms of what tools to choose for their jobs. We need to understand the metrics needed to evaluate workflows, and offer recommendations based on various criteria, such as performance, throughput, ability to handle different classes of problems, feature sets, ease of use, etc.

Overall, workflow systems will play a key role in orchestrating cross-site opportunities when it comes to linking heterogeneous scientific capabilities, computational resources, and data, and these systems should be considered a major area where collaboration across the labs will be crucial. The hope is that something along the lines of a DOE-wide Scientific Data Facility (SDF) will help drive improvements in cross-site security access and tools that will enable multi-site workflow opportunities.

List of Workflow Systems in HPC Discussed

This list is intended to provide a quick reference for the list of systems that were discussed in our session. This is not meant to be a comprehensive list and only covers those workflow systems that the participants were actively using or directly familiar with.

Fireworks - (NERSC) - FireWorks workflow software, <http://pythonhosted.org/FireWorks>. DOI: 10.5281/zenodo.10196

qdo - (NERSC) - <https://bitbucket.org/berkeleylab/qdo>

Velo - (PNNL) - http://www.pnl.gov/computing/technologies/sdm_velo.stm

PSUADE - (LLNL ASC) - http://computation.llnl.gov/casc/uncertainty_quantification

CRAM - (LLNL) - <http://spscicomp.org/wordpress/wp-content/uploads/2014/05/Musselman-High-Throughput-for-UQ-revised.pdf>

UQ Pipeline - (LLNL) - <http://computation.llnl.gov/system/files/downloads/Computation-Annual-Report-2013-LLNL.pdf> (page 21)

Kepler - (OLCF) - <https://kepler-project.org>

Swift - (ALCF, OLCF) - <http://swift-lang.org/main/>

Hadoop - (OLCF, NERSC, SNL ASC, LLNL ASC) - <http://hadoop.apache.org>

Summary

The Workflows session found that data-driven science fields are continuing to move toward the use of automated workflow systems on traditional HPC compute systems. These workflow systems typically exist to work around limitations in existing batch scheduling systems and enable coupling a series of independent codes to operate on data as it flows through the workflow session. It will be critical for HPC facility centers and HPC vendors to adapt to these workflow systems and enable them to function correctly on new systems.

Data Transfer

Session D2SD

Authors

Andrew Cherry (ALCF), Eli Dart (ESnet)

Contributors

Curt Canada (LANL ASC), Bill Collins (SNL ASC), Chris Fuson (OLCF), Mark Gary (LLNL ASC), Cory Lueninghoener (LANL ASC), David Smith (LLNL), Joel Stevenson (SNL ASC)

Executive Summary

In this session, we identified a number of best practices for managing wide area data transfer at DOE facilities. The use of dedicated data transfer nodes (DTNs) provides a central point for capacity management and monitoring, along with dedicated capacity. Deployment of Globus as our primary DTN toolset has increased accessibility of GridFTP to end users, allowing more to take advantage of its performance gains compared to more familiar tools such as scp/sftp. Collection of data transfer and sharing requirements from projects at inception will allow us to better meet oncoming capacity needs.

But many challenges remain. Increasing backbone bandwidth is only the first step—making full use of this capacity requires careful attention to factors such as DTN-to-filesystem performance, local network topologies, tools performance, and data set composition. TCP loss behavior and network security appliances can have a detrimental effect on performance and must be accounted for. Emerging needs for transparent sharing of data and integration with cloud services bring their own new set of concerns relating to security, identity management, and data integrity.

Opportunities for collaboration between DOE sites are numerous. Deployment and regular testing of our toolsets between sites can improve interoperability and identify gaps that need to be addressed, and sharing of DTN configuration and best practices will allow improved interoperability between sites. Regular meeting of working groups consisting of representatives from multiple labs is one way to encourage this.

Key Points

- DTNs are key for moving data in and out of DOE HPC centers.
- Most of the DOE HPC sites now have connectivity to the 100 Gigabit ESnet network, but some are still working on getting the full bandwidth to their DTNs.
- To fully exploit increasing backbone bandwidth requires careful attention to DTN-to-filesystem performance, local network topologies, tools performance, data set composition

and other factors. Data ingest and export from non-DOE sites such as smaller universities and businesses remains a challenge.

- The need to share and publicize data with external collaborators is a rapidly emerging need that requires careful attention. Data sharing carries concerns about identity management and security practice.
- Balancing security, which introduces overhead, with increasing performance needs is an ongoing challenge.
- Collection of data transfer and sharing requirements from projects at inception is needed to allow us to meet oncoming capacity needs.
- Network capacity needs to be able to cover the need for both heavy, sustained network transfers and large spikes over relatively short periods.

Opportunities for Collaborations Among DOE HPC Sites

- We believe that multi-facility collaboration will play a key role in our strategy over the coming years.
- HPC centers should agree upon, provide, and regularly test a common tool set and use consistent tuning practices. Sites should develop a set of common best practices and policies.
- Sharing of best practices will allow improved interoperability between sites. Regular meeting of working groups consisting of representatives from multiple labs is one way to encourage this.
- Data sharing (especially in-place sharing) is another upcoming need where a solid strategy is required. It is important that this need be addressed in a way that provides adequate safeguards for security, accountability, and control of potentially sensitive data.

Introduction

This report summarizes the findings of the Data Transfer breakout session at the 2014 HPCOR, held in conjunction with the Joint Facilities User Forum on Data-Intensive Computing. The purpose of this session was to examine current best practices and future opportunities for site-to-site data transfer, with the goal of enabling data-driven scientific discovery at DOE facilities.

During this session we sought to answer the following broad questions:

- What systems are currently in place for wide area data transfer?
- What WAN access is in place and what is needed?
- What are some best practices that you think are effective?
- How can we forecast future needs, and what changes might be needed to meet them?
- What are our biggest challenges and gaps over the coming 5-10 years?
- What opportunities exist for collaboration between DOE facilities?

Session Findings

Current State of Data Transfer

For network-based data transfer, most of the participants reported using GridFTP/Globus and bscp for large data sets. Transfers are usually initiated directly by users without staff intervention, and it was generally agreed that adoption of Globus has been quite helpful in this regard, since it makes the process considerably easier for end users. Those sites that are not yet participating in Globus reported that they are working toward adoption. Some sites have begun initial exploration of Globus Sharing (a facility that allows in-place sharing of data with external collaborators), but this is currently in the very early exploratory stages due to security and policy concerns.

Data transfers initiated by center staff are primarily used for disaster recovery and other operational functions. However, there are still cases in which user data is shipped physically on portable disks or tapes for manual ingest by site admins. At some locations this is done on a regular basis. One of the major drivers for increased network capacity is to reduce the need for these methods, which are labor intensive, slow, and vulnerable to shipping damage.

While it is usually not thought of as a protocol for large data transfers, HTTP/HTTPS is often used for data that is transferred via web portals. Another type of data transfer that is often overlooked when considering “data transfer” as a category is data transfer that does not involve POSIX file movement, such as real-time data streaming.

WAN Access

Deployment of 100 Gigabit network connectivity to the ESnet backbone is ongoing at DOE facilities. Most of the sites participating in this session now have 100Gbit to ESnet, but some are still working on getting the full bandwidth to their DTNs, and some are in pre-production testing. In some cases, alternate paths or dedicated circuits exist between DOE sites to meet capacity needs (such as DisCom at LLNL/LANL/SNL). DTNs commonly have 10Gbit network interfaces, and those that are still using 1Gbit are moving forward with planning and procuring 10Gbit equipment.

Public-facing DTNs are available at several of the surveyed sites. In more compartmentalized environments, there are internal DTNs for use with specific projects and/or specific site-to-site paths (DisCom for instance), with external/open DTNs in planning stages.

DTN pool size at participants’ sites ranges from four to 12 nodes. Increasing the number of DTNs can improve performance at some sites and will enable better use of the 100Gbit backbone. Bandwidth management (e.g. QoS) may also help address contention between data transfer traffic and other traffic in cases where dedicated circuits for data transfer are not available.

The need for network capacity is driven by a diverse set of requirements. Some experiments and user facilities (such as light sources and telescopic sky surveys) generate very large quantities of data that require ongoing transfer to other sites for processing. This necessitates a consistent level of

service over a long time period. Other usage occurs in large spikes over a short time period, particularly in the case of data preservation when project allocations end at a given facility, requiring data migration. Network capacity needs to be able to cover both of these use cases effectively.

Network security is also a significant factor, as network security appliances can add significant overhead and limit performance. Balancing security compartmentalization requirements with increasing performance needs is an ongoing challenge.

Best Practices

We agreed that the most important “best practice” for data ingest/export is to deploy dedicated DTNs with a common toolset and consistent tuning practices. DTNs provide a single place to steer users for their data transfer needs, a single point of entry for network provisioning and metrics, and a single place to deploy tools. The ESnet FasterData Knowledge Base (<http://fasterdata.es.net/>) details a number of best practices for network and host configuration and tuning, geared specifically toward DTN flows.

Globus has been adopted as the toolset of choice for DTN deployments. The GridFTP backend and its support for parallel streams makes better use of available bandwidth than serialized transfer tools, and the Globus web-based frontend has greatly increased user participation due to ease of use.

We have also found that strategically placed perfSONAR hosts are a very useful addition for troubleshooting site-to-site network performance issues.

Major Strategies and Initiatives (5-10 years)

Development and deployment of a common capability toolset across all DOE facilities is a major initiative and is essential to interoperability between sites. Development of common best practices and policies is an important component of this, and multi-facility collaboration will help achieve this goal.

Data sharing (especially in-place sharing) is another upcoming need in which a solid strategy is required. It is important that this need be addressed in a way that provides adequate safeguards for security, accountability, and control of potentially sensitive data.

IPv6 adoption is also an important initiative, if only to maintain interoperability with sites that have adopted IPv6 exclusively.

We believe that multi-facility collaboration will play a key role in our strategy over the coming years.

Forecasting of Future Requirements

There are a number of formal programmatic vehicles used to forecast requirements:

- ESnet requirements reviews:
<http://www.es.net/about/science-requirements>
- NERSC requirements reviews:
<https://www.nersc.gov/science/hpc-requirements-reviews>
- LANL ASC has an internal review every three years for large-scale planning

Another method used to forecast requirements is to track general trends in data movement and network usage. This may be done at the network level by monitoring data flows and at the host level by gathering and mining transfer log data.

For individual allocations and projects, we have found that data transfer requirements are sometimes overlooked when considering project needs. Project allocations are often focused around CPU hours, memory requirements, and storage usage estimates, with less attention given to data movement. Explicitly requesting that PIs estimate their anticipated data movement needs ahead of time would allow sites to be more proactive in ensuring these needs are met.

Challenges and Gaps

Data ingest and export from non-DOE sites such as smaller universities and businesses remains a challenge. Such sites may not have the resources for high-throughput connectivity and may end up resorting to physical shipment of media to transfer large quantities of data. Upgrades due to the NSF CC-NIE and CC*IIE programs are helping significantly in this space, but there is still more work to be done.

Numerous factors make achieving high performance over the WAN a challenge. Security enclave partitioning requires a lot of cycles for network staff, and network security appliances such as firewalls have a major performance impact. Many LAN-based protocols for data movement (HPSS mover protocol, commercial backup solutions, etc.) do not adapt well to the WAN.

The need to share and publicize data with external collaborators is a rapidly emerging trend that requires careful attention. Data sharing carries concerns about identity management and security practices, particularly when it involves the use of third-party services. These considerations are especially important on storage systems that may contain sensitive or export-controlled data. Part and parcel with data sharing is the additional challenge of managing metadata and providing data location services; shared data is only useful if collaborators can easily find the data they need. When this extends to providing public access to open data, WAN capacity planning becomes a major concern.

Other challenges we identified include tuning DTN pools to make the best use of parallel file systems (e.g. GPFS) and integrating existing tools with cloud-based storage and computing systems.

We believe the main gap we are currently facing is the discrepancy between available backbone network capacity and our current ability to use it. Security considerations, variances in data set composition, limits in DTN pool size, contention and restrictions on local networks, and limitations in existing toolsets are among the many factors that contribute to this.

Opportunities for Collaboration

Because data transfer by its nature involves interaction between multiple sites, there are many opportunities for collaboration between DOE facilities. Sites can assist each other in developing common toolsets and standard best practices. Regular testing of tools and capabilities with broad sharing of results can help identify where gaps remain. A good example of this is the ALICE project at CERN (see the Network Traffic section at <http://alimonitor.cern.ch/>). There are also opportunities for increased collaboration between DOE sites and the Globus project, such as helping Globus make better use of large DTN pools that mount parallel file systems.

Even within a given organization, collaboration is critical for success. Data transfer involves many potential bottlenecks between the source and destination: storage hardware, file systems, data transfer nodes, local networks, security appliances, wide area networks, and so on. This multi-disciplinary nature makes it critical to build cross-functional collaborations between admins in all of these disciplines.

Attendees

U.S. Department of Energy Office of Science (DOE SC)

Laura Biven	Senior Science and Technology Advisor Office of the Deputy Director for Science Programs DOE Office of Science
-------------	--

Office of Advanced Scientific Computing Research (ASCR)

David Goodwin	NERSC Program Manager
Barbara Helland	Facilities Division Director
Lucy Nowell	Data and Visualization

Argonne National Laboratory (ANL)

Argonne Leadership Computing Facility (ALCF)

William Allcock	Kevin Harms
Susan Coghlan	Michael Papka
Andrew Cherry	Venkat Vishwanath
Richard Coffey	Yao Zhang

Los Alamos National Laboratory (LANL)

Advanced Simulation and Computing (ASC)

Curt Canada	Kaki Kelly
Kyle Lamb	Dee Magnoni
Cory Lueninghoener	Laura Monroe
Bob Kares	

Lawrence Livermore National Laboratory (LLNL)

Advanced Simulation and Computing (ASC)

Sasha Ames	Robin Goldstone
Blaise Barney	Ming Jiang
Jeff Cunningham	Jeff Long
Tim Fahey	Jim Silva
Mark Gary	David Smith

Lawrence Berkeley National Laboratory (LBNL) (“Berkeley Lab”)

Energy Sciences Network (ESnet)

Eli Dart

National Energy Research Scientific Computing Center (NERSC)

Katie Antypas
Wes Bethel
Jeff Broughton
Shreyas Cholia
Chris Daly
Sudip Dosanjh
Brent Draney

Kjiersten Fagnan
Richard Gerber
Jason Hick
Prabhat
David Skinner
Nick Wright

Physics Division

Stephen Bailey

Oak Ridge National Laboratory (ORNL)

Oak Ridge Leadership Computing Facility (OLCF)

Ashley Barker
Fernanda Foertter
Chris Fuson
Norbert Podhorszki

Doug Fuller
John Harney
Clay England
Julia White

Spallation Neutron Source (SNS)

Stuart Campbell

Sandia Advanced Simulation and Computing (SNL ASC)

Bob Ballance
Chris Beggio
Bill Collins
Rudy Garcia
Karen Haskell
David Karelitz

John Noe
Dino Pavlakos
Joel Stevenson
Craig Ulmer
Andy Wilson

San Diego Supercomputing Center (SDSC)

Glenn Lockwood

National Center for Atmospheric Research (NCAR)

Ilana Stern

DISCLAIMER

This document was prepared as an account of work sponsored by the United States Government. While this document is believed to contain correct information, neither the United States Government nor any agency thereof, nor the Regents of the University of California, nor any of their employees, makes any warranty, express or implied, or assumes any legal 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 its trade name, trademark, manufacturer, or otherwise, does not necessarily constitute or imply its endorsement, recommendation, or favoring by the United States Government or any agency thereof, or the Regents of the University of California. The views and opinions of authors expressed herein do not necessarily state or reflect those of the United States Government or any agency thereof or the Regents of the University of California.