

Workshop Report:
2022 International Symposium of Quantitative Codesign of Supercomputers
Version: February 22, 2023



Acknowledgements

We would like to acknowledge the efforts of the following people who not only made our symposium possible, but worked to make our symposium outstanding: SC'22 workshop committee chair Rio Yokota, vice chair Tjerk Straatsma; the anonymous workshop proposal reviewers who provided helpful guidance; SC'22 planning committee conference vice chair William Kramer; virtual logistics chair Lori Diachin; our web designer Douglas Edwardson; and the audio-visual team that provided support for our symposium.

Thank You!

TABLE OF CONTENTS

1.	BACKGROUND ON QUANTITATIVE CODESIGN	3
2.	PURPOSE OF THE WORKSHOP	4
3.	WORKSHOP STRUCTURE	6
3.1	AGENDA.....	6
3.2	A HYBRID FORMAT: ACCOMMODATING IN-PERSON AND REMOTE PARTICIPATION	6
4.	WORKSHOP OUTCOMES.....	8
4.1	HPC CONTRIBUTIONS	8
4.2	WORKSHOP FINDINGS.....	9
5.	POST-WORKSHOP RECOMMENDATIONS AND “NEXT STEP” STRATEGIES	10
5.1	CONTINUE THE WEBSITE (LINK)	10
5.2	DISSEMINATE THE WORKSHOP REPORT	10
5.3	DISSEMINATE THE POSITION PAPERS	10
5.4	TRACK POTENTIAL MISSION FURTHERING OPPORTUNITIES.....	10
5.5	ADVANCE THE QUANTITATIVE CODESIGN AGENDA WITH A 2023 SYMPOSIUM	10
	APPENDIX 1 – RELATED ACTIVITIES	11
	APPENDIX 2 – SPEAKER BIOGRAPHIES	12
	APPENDIX 3 – ORGANIZING COMMITTEE AND PROGRAM COMMITTEE	15
	APPENDIX 4 – ATTENDEES & WORKSHOP PHOTOGRAPHS.....	16
	APPENDIX 5 – POSITION PAPERS.....	18
	POSITION PAPER #1: THOMAS JAKOBSCHKE	19
	POSITION PAPER #2: BRANDON KAMMERDIENER.....	23
	POSITION PAPER #3: JIM BRANDT	26
	INVITED POSITION PRESENTATION: WILLIAM KRAMER.....	29

2022 International Symposium on the Quantitative Design of Supercomputers
Held in conjunction with Supercomputing '22
Dallas, TX – November 13, 2022

1. Background on Quantitative Codesign

The Quantitative Codesign of Supercomputers symposium is an annual workshop series that aims to significantly improve the effectiveness of high-performance computing through bringing about increased understanding of current limitations and improved development processes. This symposium considers combining two methodologies—collaborative codesign and data-driven analysis—to realize the full potential of supercomputing. For full potential of supercomputing we consider everything pertaining to output production, including but not limited to the performance of applications, system software, workflows, health of hardware. Our centers store vast sums of information, yet using this data is a demanding task. To a large extent the difficulty in obtaining quantitative insight has to do with discovering, accessing, and analyzing the right data. Codesign also presents formidable challenges, e.g. on how to use the data collected on current systems to facilitate the (potentially very different) design of next-generation supercomputers and successfully support our upcoming environments. Quantitative codesign offers a collaborative evidence-based approach to address our existing needs and our upcoming ambitions. This symposium was created to bring together leaders in the field to review current efforts across centers and discuss areas that show potential.

Over the past decade, there has been a growing awareness of the multi-faceted benefits we can derive from data-driven strategies like Quantitative Codesign. This increasing awareness, along with improvements in Machine Learning (ML) technologies, have driven vendors, operations staff, and application developers to espouse integrating an ever-increasing level of instrumentation into their products. The time is ripe for turning this vast trove of available information and the incredible advances in analysis technologies it represents into appropriate knowledge and understanding. Doing so would create a feedback loop that could assist vendors and software developers in their designs. The recent National Strategic Computing Initiative Update Report has recommended that we promote timely access for developers of technologies, architectures, and systems to carry out the research needed to create the future computing software ecosystem, and Quantitative Codesign provides a solution to the ‘access problem’ of these extremely rare machines. If the future envisioned by the CSESSP report is to be realized, our software base will require significant investment in both modified and new code — an activity enormously assisted by Quantitative Codesign. There is no disagreement that more knowledge is good though there is still lack of concurrence across HPC stakeholders as to the cost/benefit tradeoff for varying fidelities of information collection and long term storage. The benefits of Quantitative Codesign will come through integrating design processes with more detailed knowledge of the interactions of the various components within the HPC ecosystem.

Quantitative Codesign is also essential for addressing challenges brought about by the recent trend of increasing heterogeneity and varied accelerators in HPC architectures. For example, many HPC machines now incorporate alternative types of memory alongside conventional DDR SDRAM. Technologies such as

"on-package" or "die-stacked" DRAM as well as non-volatile RAMs can provide distinct advantages compared to conventional DRAM, including higher performance as well as cheaper and more energy efficient storage per byte. Each of these technologies also comes with its own limitations, such as smaller capacity or less bandwidth for reads and writes. Further complications arise because some of these new technologies can interface directly with processor caches, while others can only be accessed through peripheral devices, such as GPUs or other accelerators.

Quantitative Codesign could mitigate many of the current problems with allocating and managing such heterogeneous resources effectively. Detailed knowledge of application demands will enable architects to make better decisions about how to select and organize computing and memory hardware. This approach can also help system software, including operating systems, compilers, and runtime software, distribute the available hardware resources among applications more effectively. Codesigned system software could utilize knowledge from new data sources for better energy efficiency and workflow management. Integrating high-level profiling and analysis with low-level resource management routines will enable these systems to implement new policies that respond flexibly to changes in application demands and could potentially expose important new efficiencies on platforms with heterogeneous hardware.

2. Purpose of the Workshop

The purpose of the workshop was to build the necessary community support to build up and foster concrete implementations of quantitative codesign. As architectural options expand in type and complexity, the need for a quantitative basis to drive architectural directions becomes increasingly urgent. We do not have the primary mission to raise awareness of an individual's research; rather we wish to bring more wide-ranging interactions highlighting vision and positions and stimulating discussions.

Any shortfall in our detailed understanding of operations and performance impacts the whole spectrum of stakeholders. Whether providing hardware architectures, system software, application programming environments, or production run-time environments, having the appropriate knowledge to optimize the interaction and configuration of all of these critical components as well as the evolution of the HPC ecosystem is critical to continued growth. The rapidly changing HPC landscape demands a codesign that effectively uses the data collected on previous and current systems to facilitate the design of next-generation supercomputers and successfully support our upcoming environments. Specifically, we would like to bring increased clarity for our challenges and opportunities.

- **Challenges:** We have important issues to resolve, but we are not starting from scratch. HPC computing centers already collect a wealth of information on the health, usage, and efficiency of our machines, workflows and programming environments. While collection and analysis of this information has evolved and improved over the years, there are still severe gaps that have left us unable to provide the knowledge that is needed by hardware and software vendors, system operations staff, application developers, and user groups to create and operate highly efficient and secure large scale HPC systems. Would-be users of this information face difficulties in obtain insight from the collected data a timely manner, and efforts to provide both data and analysis means are currently fragmented across centers both at national and international levels. The infrastructure to collect, store, share and analyze the volumes of available information is a core capability—yet, many barriers remain due in large part to the many stakeholders and insufficient coordination, but also due to data privacy and security issues. With many new potential information sources in future systems, we must quickly identify and address critical requirements and gaps across the various stakeholders. Doing so will enable us to create collective and collaborative solutions that address both existing challenges and emerging needs and effectively support our upcoming HPC environments. The nature of this challenge suggests that it is an excellent opportunity for a codesign approach. Codesign is defined as the process of jointly designing interoperating components of a computer system—in particular: applications, algorithms, programming models, system software, as well as the hardware on which they run, and the facilities hosting them. Designing solutions based on

intelligence derived from the data collection and analysis processes described above are henceforth referred to as Quantitative Codesign of Supercomputers.

Making progress at the highest end of HPC without access to the needed data can be compared to being asked to fly an airplane at night without sufficient instrumentation. Vendors are provided with example applications to target, but often lack a true understanding of where inefficiencies manifest on full scale workloads. Furthermore, computer architecture simulators face an inevitable challenge in trying to incorporate all the critical performance-killing attributes of current generation technologies and their integration: a simulation that includes all details of the architecture, from the chip micro-architecture up to infrastructure, would take forever to run. For this reason, simulations must make tradeoffs between the accuracy of their representation and the required modelling time. Hence the vendors miss opportunities for improvement. Moreover, users often only have feedback on operating efficiency at the granularity of total application execution time. Low-level interactions frequently cause substantial performance degradations that users are unable to explain. Likewise, operations staff often lack knowledge of application resource utilization and cannot diagnose the longer run times experienced by the users. In addition, operations staff cannot ensure secure operations without an understanding of normal (expected) behavior and anomalies that deviate from that. Since root causes go undiagnosed on current systems, next generation systems will also fail to address the very same problems.

- **Opportunities:** First and foremost, we wish to discuss the merits of a coordinated effort to bring together the helpful data from each stakeholder in the codesign space into a framework where data discovery and access is straightforward regardless of data source while respecting data privacy and security concerns. The envisioned Quantitative Codesign environment would pull together data traditionally held by disjointed communities (e.g., sysadmins, application teams, vendors, and so on) into a framework where the needed data is easily accessible. This framework would provide flexible but secure mechanisms for data providers who wish to share their data with others including application teams, vendors, facilities, operations, and system software researchers. In many cases, we seek to bring together data that is currently being produced although not generally known or utilized for a variety of reasons; in a few instances, we seek to extend and provide new data collection capabilities.

For example, one area that is ripe for integration with Quantitative Codesign processes is the intersection of application development and run-time environments. In the past few years Continuous Integration (CI) has been widely adopted by development teams to continuously test development efforts. As part of these CI efforts, developers test across a variety of platforms on a daily basis and typically provide a pass/fail result for each. Introducing targeted run-time data collection (e.g., memory, application & hardware counters, MPI, OpenMP, GPGPU, I/O, energy consumption) and quantitative analysis into this process would enable feedback to users and identify issues within applications, compiler capabilities, runtimes, and differences across platform architectures that ultimately would drive improvements across the spectrum of stakeholders.

Integrating Quantitative Codesign capabilities with existing design processes will enable more effective solutions across the computing stack. Information derived from monitoring and analysis would provide valuable insight for users, application developers, system architects, and facility designers as to how, and why, applications make use of the underlying system resources. Furthermore, by identifying the appropriate stakeholders and introducing them to information originating from diverse collection regimes, this symposium seeks to facilitate the discovery and sharing of potentially useful intelligence among larger teams and communities. In doing so, this approach also has the potential to spark further discussions and research on how to collect, employ and share this information more effectively. Thus, there is significant opportunity for discoveries that will not only increase application performance, but also benefit the broader HPC and scientific communities.

3. Workshop Structure

The Quantitative Codesign of Supercomputers symposium took place during the opening day of the 2022 Supercomputing conference. Due to COVID conditions at the time, the workshop was held in hybrid model with both in-person and virtual attendees and speakers. The workshop was framed in the Symposium format to achieve the kind of deep interactions that lead to change within HPC. Our preference for audience interaction was in response to the state of the field (which we see as in its infancy).

3.1 Agenda

Given our desire to bring more wide-ranging interactions highlighting vision and positions and stimulating discussions, we developed a schedule designed to facilitate these interactions (see Table 1 below). In particular:

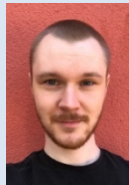
- The keynote speaker was chosen based on his long history in HPC with work that spans all areas of codesign including novel architectures, system and application software, tool development, performance diagnostics and more, in both lab and academic environments.
- Three speakers were chosen who, as an aggregate, provided codesign perspectives on common misunderstandings of what an ISA actually provides, ways in which AI can be employed in this field, and experiences in bringing about a holistic monitoring system at a major supercomputer center.
- A collection of position papers from an international collection of experts with diverse backgrounds in codesign, HPC system software and middleware research, center wide monitoring and operational aspects, bringing HPC products to market, and application / libraries.
- A moderated discussion of audience, speakers, and panelists was included to enable both technical discussions and community-building.

3.2 A Hybrid Format: Accommodating In-Person and Remote Participation

As with our previous Symposium, the COVID pandemic had an impact on the format and character of the workshop. This was the second time for the SC series of conferences to ever have a hybrid format: SC22 supported both in person attendees at the Dallas Convention Center in Dallas, Texas and remote attendees through the revamped SC22 online platform, Zoom and Sli.do. The role of the session chair and organizer remained largely the same as in previous years with some adjustments and increased responsibilities to account for remote participation by speakers and attendees. The Quantitative Codesign of Supercomputers symposium was presented via *Live stream sessions*. Under this format, content was recorded by AV technicians at the convention center and sent to remote participants in real time via Vimeo. Remote presenters connected via zoom (see Figure 1). For all remote symposium presenters, we arranged for an internet assessment on the day of the symposium prior to the symposium start. This was used to ensure no fallback measures were needed. All remote speakers were able to participate as planned.

The Symposium's program committee considered SC'22's Live Stream hybrid format a "mixed-bag". Last year, SC'21 utilized a new *Hubb virtual interface* to provide coordination of in-person presentations working in concert with virtual zoom interface; last year's overall SC'21 experience was positive given Hubb was unfamiliar. For SC'22, it was decided to emphasize the in-person (in Dallas) experience. The new system was sufficiently complicated to require a fair amount of AV knowledge and/or training. Unfortunately for our symposium (which was scheduled on Sunday morning, the first timeslot of the whole conference), our AV support person was unprepared and we had significant AV issues during the first half of our symposium: (1) remote people coming in through the SC'22 website did not have audio; (2) our two remote speakers had difficulties starting their zoom presentation; (3) our in-person audience experienced a 20 minute delay while the AV person tried to figure out the AV set-up; (4) we frequently had serious feedback issues during the course of the day. It is our belief that these problems were largely a result of insufficient training for the AV person, and that the same technology could provide a positive experience next year in SC'23.

Table 1 – Symposium Agenda

Duration	Speaker/Panelist	Abstract
5 mins	Terry Jones	Welcome
45 mins	 Jose Moreira	Predictions are hard, especially about the future The lifetime of a leadership computing system spans about 10 years: It takes 5 years from conception to delivery, and then the systems stays in production for an additional 5 years. A lot can change in 10 years: Strategic priorities, policies, technologies, usage patterns, economic landscape. Yet, we have to make this work. In this talk, I will discuss some of the approaches we can use to design a high-performance computing system that will be relevant 10 years in the future. I will also discuss why sometimes we have to look beyond what we actually know, in order to produce true game-changing systems.
20 mins	 Ronald Minnich	Wait! This is the Bad Place! – common misconceptions of hardware and impact on codesign Designs of integrated stacks on, e.g., x86 systems, often build on a fiction that does not exist: that of a 1994-era Pentium over which the software has complete control. This is a fiction: in fact, kernels run on a virtual system over which they have little control. This would be fine if it did not affect performance, but it can in the end have significant throughput impacts, e.g., on some modern systems, entire sockets can stop for 1/2 a second at a time. I will discuss some of the cases in which these Potemkin Villages have caused real trouble, and suggest possible ways to deal with them on old (x86) and new (RISC-V) systems.
20 mins	 Devesh Tiwari	Learning to unlearning (some) conventional wisdom in HPC system design and operation Traditionally, we have assumed that HPC users are fairly boring and that their workloads often do similar things repetitively. Their "boring" nature has served us well so far -- we could design "boring" systems and get away with it. But, now things are changing and changing fast. Our HPC workloads and users are becoming interesting and, often, are surprising us with new trends and behavior. That means it is springing excitement into our lives. We need to design interesting solutions, and come out of our boredom.
15 mins	 William Kramer	A Pathway to Achieve the Holy Grail for Efficient Quantitative Co-Design Systems Today, many supercomputing organizations perform system evaluation and analysis, often utilizing different data and tools. But within sight is now the ability to use high fidelity, continuously collected system wide data collection, combined with vendor independent, community SW tools to do real-time system management, application performance improvements all the way to long term Quantitative Co-Design. When merged with models of the next generation of applications and methods, we may be able to rapidly and fully evaluate many configuration and systems to optimize the next generation technologies.
15 mins	 Thomas Jakobsche	Challenges and Opportunities of Machine Learning for Monitoring and Operational Data Analytics in Quantitative Codesign of Supercomputers This work examines the challenges and opportunities of Machine Learning (ML) for Monitoring and Operational Data Analytics (MODA) in the context of Quantitative Codesign of Supercomputers (QCS). MODA is employed to gain insights into the behavior of current High Performance Computing (HPC) systems to improve system efficiency, performance, and reliability (e.g. through optimizing cooling infrastructure, job scheduling, and application parameter tuning).
15 mins	 Brandon Kammerdiener	A Quantitative Approach for Guiding Data Management on Complex Memory Architectures Proliferation of real-time and AI-driven decision making continues to fuel the need for ever faster access to larger sets of data in memory. At the same time, increasing demands for high-density sharing are leading to more complex memory architectures with rich opportunities for addressing the diverse needs of applications under various cost, performance, and power constraints. We propose to address these challenges through a quantitative approach that leverages detailed profiling and analysis of application behavior to steer data management of complex memory platforms.
15 mins	 Jim Brandt	AppSysFusion: CoMingling of appropriate data to drive Codesign of Applications, HPC Platforms, and Monitoring, Analysis, and Feedback Infrastructure The goal of building HPC systems is to enable execution of large-scale user application workflows in an efficient and performant manner. Performance here is multi-dimensional and includes not just a particular application's time-to-solution, but the aggregate throughput of all applications submitted (workload) and energy spent. The aggregate HPC system power draw must always remain within a contract envelope.
30 mins	Terry Jones	Community Building Discussion with Audience and Closing Remarks

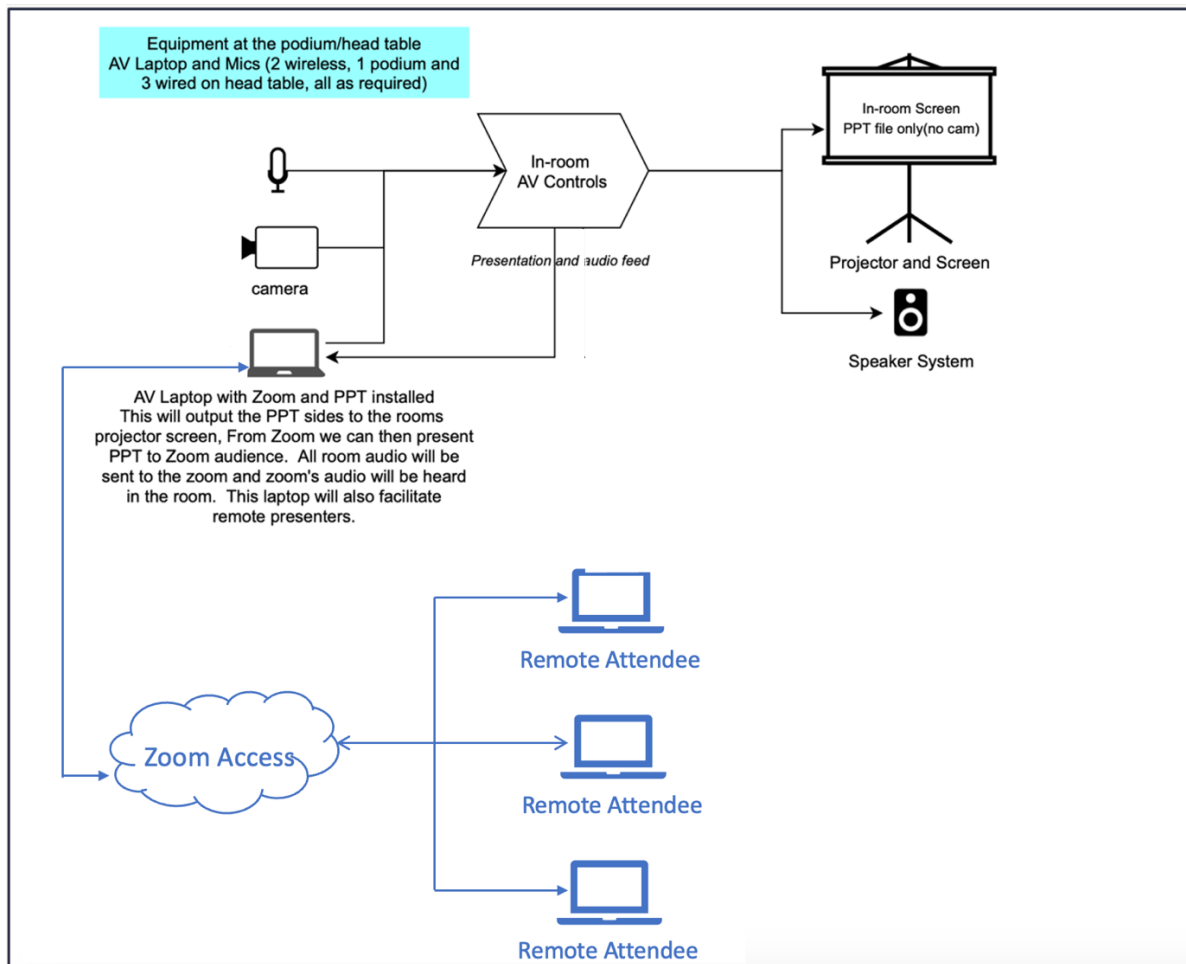


Figure 1 – Logistical Setup of Live Stream Format Used By Symposium.

4. Workshop Outcomes

4.1 HPC Contributions

The following positive results have been achieved with from the workshop:

- A large group of high performance computing professionals came together to pursue community building
- Monitoring journals (outcome and strategy) were discussed and templates provided to guide the process of data collection and the use of these data
- Videos of the invited talks and panels were recorded by SC's Live Stream AV team
- Discussion on Vision and Possibilities of Quantitative Codesign of Supercomputers were discussed, and ideas for future work were identified
- This workshop report was written to document the results

In addition, monitoring journals (outcome and strategy) were discussed and templates provided to guide the process of data collection and the use of this data.

4.2 Workshop Findings

Consensus Views – The following points were generally agreed upon by those attending:

- Recent computer architectures intended for high performance computing are becoming increasingly heterogeneous and are growing in complexity.
- Likewise, the workloads we run on them (e.g., the growth of machine learning) and the performance metrics we care about (e.g., the new concern over electrical power usage) is changing.
- Recent HPC systems have taken 5 years to design, then are in production for around 5 years.
- Perhaps the most nebulous aspect of evaluating a computer system under design is the selection of workloads. Note that the 5 year design cycle for computer architectures makes talking about their workload especially challenging since workloads are constantly changing. Further, workloads desire to take advantage of new hardware capabilities and therefore wish to adapt themselves.
- There's a temptation to assume that we can build software stacks (including kernels and runtime systems) on a stable Instruction Set Architecture, or ISA, (e.g., the 1994 era Pentium ISA) without regard to impacts from low-level hypervisors. The reality is that these low-level hypervisors can have dramatic impacts and must be carefully considered.
- Smart runtime systems have shown promise in mapping applications to different complex architectures without requiring application modification (e.g., effective use of complex memory systems with multiple tiers).
- Making our monitoring data holistic and available to machine learning could have huge impact on computer centers in terms of efficiencies, problem detection and problem resolution. This is particularly true of software errors and cascading errors.
- An emerging need is optimizing for an entire workflow (as opposed to a single element or application execution).
- The idea of a “slack channel” like capability that is able to provide news for this topic area for people interested in the topic area was suggested.

Value Proposition – The following points in support of Quantitative Codesign's value proposition were documented:

- There are a number of historical assumptions (“Conventional Wisdom”) about how workloads and systems behave that we are finding are incorrect as we have more insight into the data (e.g., workloads are repetitive). This motivates further exploration of data and opportunities to affect future design.
- A smart runtime system, ECP's SICM project, demonstrated 1.4x to 7x improvements in performance by monitoring the usage patterns of an HPC application (WarpX, Nalu Wind, Coral proxy apps) and guiding resource management on a complex computer architecture.
- A holistic monitoring system, NCSA's Blue Waters monitoring system, has proven to be an effective tool for uncovering the root cause of application performance variability.
- Codesign was shown to be an effective and efficient way to design successful machines (e.g., IBM's Blue Gene machine, Riken's Fugaku machine).

Open Questions and Challenges – The following items have no acceptable solutions at present or are considered unresolved:

- Social challenge:
 - Computer architectures and software stacks have long timelines; we have problems finding people that are ready to engage for such long.
 - Coming from DoE R&D background, the work that we did was often in 3-years sparks and then we moved on to another problem. This is not the right approach for this topic.
 - One key aspect to get this in the right time frame is getting the right people on board.

- One answer to this is funding.
- Data Challenges:
 - Analysis paralysis: The unwillingness to make a definitive decision until that next bit of data is available.
 - Too much data is proprietary, confidential, or there are sharing concerns. What role does Anonymisation play?
 - Lack of ownership and leadership in this topic
 - Making data FAIR (Findable, Accessible, Interoperable, and Reproducible)
 - Overhead associated to monitoring
 - Budget: where is there money to do this?
- Characterizing Our Workload
 - Existing Workloads: In core design not too much concern about workloads, but we do not have a clear decision on workloads that are used for AI and codesign.
 - Role of ML for optimization: HPC is very broad. We need tools for computational scientists not so used to HPC. But ML tools get often stuck in local minimal and are not generalizable. If you get from a platform to another you might get a different answer.
- Funding:
 - Next generation projects after Exascale need a clear investment in this area.

Note: We are currently preparing a transcript that reflects our notes from dialogue during the workshop. This transcript will be added in an updated version of the report when it is ready.

5. Post-Workshop Recommendations and “Next step” Strategies

There are a number of recommended “next steps” that should be followed to increase the usability of quantitative codesign of supercomputers.

5.1 Continue the website ([Link](#))

Provide ongoing support to Quantitative Codesign of Supercomputers website. This web presence becomes an anchor for announcements and a source to discover resources and pertinent email addresses.

5.2 Disseminate the Workshop Report

Providing this post-workshop report of the event will be an important resource for the symposium’s community building objective. The contact data of the participants interested on receiving the report have been collected and will be used to spread the report in the community

5.3 Disseminate the Position Papers

Providing this post-workshop access to the 4 position papers.

5.4 Track Potential Mission furthering opportunities

This follow-up activity is to ensure that a wide segment of high performance computing is monitored for events, interactions and publications for opportunities to advance high performance computing through quantitative codesign concepts.

5.5 Advance the Quantitative Codesign agenda with a 2023 Symposium

Finally, we are encouraged to repeat the workshop in 2023. This third workshop should consider emphasizing how quantitative codesign can assist in system software development. We would include middleware and runtime systems in our scope. In post discussions of the organizing and program committee, there was a consensus to propose the 360 minute format at SC’23, which would allow additional engagement of the community through position papers and extension of the discussion period, while stipulating that we could fall back to the 180 minute format if required.

Appendix 1 – Related Activities

Among the related activities that we wish to augment are the following:

- The Center and Application Monitoring Session held during the ECP Annual Meeting.
- The [International Workshop on Monitoring and Operational Data Analytics \(MODA\)](#) held with the annual ISC High Performance conference.
- The [Workshop on Monitoring and Analysis for High Performance Computing Systems Plus Applications](#) (HPCMASPA) held with the annual IEEE Cluster conference.
- The Workshop on Performance Monitoring and Analysis of Cluster Systems (PMACS) held with the annual Euro-Par conference.

Each of these related activities share an interest in the wealth of information exposed by these systems about how the system resources are being utilized. Our Symposium is unique in its emphasis on applying data to improve the codesign process. The Quantitative Codesign Symposium also has a distinguishing format and venue.

Appendix 2 – Speaker Biographies

Terry Jones

Terry Jones is a Senior Research Staff member at Oak Ridge National Laboratory (ORNL) where he has worked since 2008 in the Computer Science and Mathematics Division (CSMD) as a Computer Scientist. Prior to that, he held a Computer Scientist position at Lawrence Livermore National Laboratory (LLNL). Terry earned a Master of Computer Science degree from Stanford University. Terry's research interests include system software for high performance computing, runtime systems and middleware, parallel and distributed architectures; performance monitoring; memory and storage systems; distributed clock synchronization, and resilience for complex distributed systems.

Jose Moreira

José E. Moreira is a Distinguished Research Staff Member at the IBM Thomas J. Watson Research Center. He received a B.S. degree in physics and B.S. and M.S. degrees in electrical engineering from the University of Sao Paulo. He received a Ph.D. degree in electrical engineering from the University of Illinois at Urbana-Champaign. Since joining IBM in 1995, Dr. Moreira has worked on a variety of high-performance systems, including two ASCI systems (Blue Pacific and White) and the Blue Gene/L supercomputer, for which he was the System Software architect. Dr. Moreira has been responsible for various architectural and micro-architectural innovations in the three most recent generations of POWER processors. He conceived the POWER10 matrix unit, the first of its kind in a commercial processor. Dr. Moreira is a Fellow of the IEEE (Institute of Electrical and Electronics Engineers) and a Distinguished Scientist of the ACM (Association for Computing Machinery).

Ronald Minnich

Ronald Minnich is a Senior Staff Software Engineer at Google. Dr. Minnich received his M.S. degree in Electrical Engineering at the University of Delaware, and his Ph.D. in Computer Science from the University of Pennsylvania. Before joining Google, he was a technical staff team lead at Los Alamos National Laboratory. He then joined Sandia National Laboratories as a distinguished member of the technical staff. Dr. Minnich has been writing firmware for 40 years, starting with the z80 and 6800. He's also a long time contributor in the Unix, BSD, Plan 9, and Linux communities. He started the LinuxBIOS project in 1999, which was renamed to coreboot in 2008 and is now used in tens of millions of Chromebooks. A more recent effort, LinuxBoot, is now part of the Linux Foundation and aims to bring the benefits of a full Linux kernel to several firmware environments, including coreboot, u-boot, and UEFI.

Devesh Tiwari

Professor Devesh Tiwari is an educator and researcher at Northeastern University where he directs the Goodwill Computing Lab. His group innovates new solutions to make large-scale classical HPC systems and quantum computing systems more efficient, reliable, and cost-effective. Before joining the Northeastern faculty, Devesh was a staff scientist at the United States Department of Energy (DOE) Oak Ridge National Laboratory. Devesh was recognized with multiple awards including the DSN Dependability Rising Star Award, the NSF CAREER Award, and the Facebook Faculty Research Award. Devesh's research group has lowered the barrier to entry and accelerated the R&D efforts in multiple emerging computer systems areas including HPC, quantum system software, serverless computing, and AI-driven data center optimizations, via open-sourcing novel software artifacts and datasets. The research contributions from his excellent PhD students have been recognized with many best paper nominations and fellowships/awards. For his teaching and mentoring contributions, he was awarded the Professor of the Year by the Northeastern University chapter of the IEEE Eta Kappa Nu honor society. Devesh has also introduced several novel peer-review elements in the computer systems community in his role as the program co-chair/track co-chair for various conferences. Most recently, he was the Technical Program

Committee Co-Chair for HPDC'22 and is the overall Technical Program Committee Co-Chair for IPDPS'23. He is an Associate Editor for Transactions of Parallel & Distributed Computing (TPDS), Transactions of Storage (ToS), and Journal of Parallel & Distributed Computing (JPDC). He was recognized with the TPDS Editorial Excellence Award for his exceptional contributions to the TPDS journal as an editor.

Bill Kramer

William T.C. Kramer is the Executive Director for the Illinois OVCRI New Frontiers Initiative and PI/ Director of the Leadership-class Blue Waters Project (<https://bluewaters.ncsa.illinois.edu>). He holds an appointment as a full Research Professor of Computer Science at Illinois. The Blue Waters system, the largest supercomputer Cray has ever built, is the first general purpose, open science, sustained-petaflop supercomputer placed into service in 2013, delivering over 35 billion core*hours of computing to date. It is the most powerful resource for the nation's open-science researchers. Blue Waters is a project with an overall cost of over \$520M to support thousands of researchers doing Frontier Science and Engineering research that is not possible any other way. Now Blue Waters is now devoted to NGA related geospatial investigations. In addition to being the Blue Waters Director, Kramer is a full Research Professor of Computer Science in the Computer Science department at UIUC and has been the PI of the NSF funded Global Initiative to Enhance @scale and distributed Computing and Analysis Technologies (GECAT) project, the DOE/ASCR funded Holistic Measurement Driven Resiliency HMDR award that studies failure and resiliency for exascale systems, several contracts with DOE laboratories and a OTA agreement with NGA/SOSSEC.

Previously, he held leadership roles as General Manager of the National Energy Research Supercomputing Center (NERSC) and was a Branch Chief in NASA's Numerical Aerodynamic Simulation division at Ames Research Center.

He is the author of over 60 papers about high end computing and data analysis and served on HPC advisory committees around the world. He is one of the founding executive directors of the Joint Laboratory for Extreme Scale Computing (JLESC) which is a collaborative organization consisting of 7 organizations in 4 countries researching extreme scale computing and data analysis. He was the General Chair of the International SC05 conference as well has many other high level volunteer leadership positions. He is the recipient of multiple awards from NASA, the Department of Energy, the Association of Computing Machinery and the Digital Computer Users Society, including one that recognized his contributions to returning the Space Shuttle to flight after the Challenger accident and another for establishing a \$400M research program for improving the effectiveness of the US air traffic control systems. He has also led the creation and was responsible for NASA's TS-SCI computing facility and holds current DOD clearances.

Blue Waters was the 20th supercomputer Kramer deployed and/or managed. He also deployed and managed large clusters of systems, several extremely large data repositories, some of the world's most intense networks and also been involved with the design, creation and commissioning of six "best of class" HPC facilities. He is known for developing the Sustained System Performance (SSP), Effective System Performance (ESP) and PERCU evaluation methods for large scale systems.

Bill holds a BS and MS in computer science from Purdue University, an ME in electrical engineering from the University of Delaware, and a PhD in computer science at UC Berkeley. Kramer's research interests include large-scale system performance evaluation, systems and resource management and scheduling, system resiliency and fault detection, large scale system monitoring and assessment and cyber protection. Bill has certifications in very large IT project management from GSA and DOE. Bill advises and consults around the world on large-scale systems and facilities and their use.

Thomas Jakobsche

Thomas Jakobsche is a research assistant in the High Performance Computing group and a Ph.D. student in the doctoral program Data Analytics at the University of Basel, Switzerland. He received his M.Sc. and B.Sc. degrees in Computer Science from the University of Basel, Switzerland. His research interests include Monitoring and Operational Data Analytics on High Performance Computing systems.

Brandon Kammerdiener

Brandon Kammerdiener is a PhD student in Computer Science at the University of Tennessee, Knoxville. During his graduate studies, he has created multiple software tools and frameworks for understanding and managing application data usage on emerging memory architectures. He is also one of the lead contributors to the Simplified Interface to Complex Memory (SICM), a DOE ECP subproject that aims to enhance data management on supercomputing platforms with complex memory hierarchies.

Jim Brandt

James (Jim) Brandt is a Distinguished Research Staff Member (Computer Scientist) at Sandia National Laboratories. Jim's research interest for the past two decades has been in holistic data-driven analysis of HPC eco-system resource utilization and state. He leads the development effort for Sandia's Lightweight Distributed Metric Service (LDMS) which has been in production use for a decade and installed on largescale systems across the DOE and NSF. Jim also leads SNL's AppSysFusion project, which enables run time combined application+system monitoring, through the interoperability of LDMS with other tools including Kokkos, Darshan, and Caliper. Jim leads work in the area of application of AI/ML to modeling and optimization of application resource utilization and anomaly detection. Jim has a M.S. degree in Computer Engineering from Santa Clara University and a B.S in Physics from California State University Hayward.

Appendix 3 – Organizing Committee and Program Committee

Workshop Organizing Committee

- Terry Jones - Oak Ridge National Laboratory, USA
- Estela Suarez- Jülich Supercomputing Centre & University of Bonn, Germany
- Ann Gentile - Sandia National Laboratories, USA
- Michael Jantz - the University of Tennessee, USA

Workshop Program Committee

- Jim Brandt - Sandia National Laboratories, USA
- Florina Ciorba - University of Basel, Switzerland
- Hal Finkel - US DOE office of Advanced Scientific Computing Research, USA
- Lin Gan - National Supercomputing Center, Wuxi, China
- Maya Gokhale - Lawrence Livermore National Laboratory, USA
- Thomas Gruber - Friedrich-Alexander-University Erlangen-Nuernberg, Germany
- Oscar Hernandez - nVidia, USA
- Jesus Labarta - Barcelona Supercomputing Center, Barcelona, Spain
- Hatem Ltaief, King Abdullah University of Science and Technology (KAUST), Saudi Arabia
- Yutong Lu - Director of National Supercomputing Center in Guangzhou, China
- Esteban Meneses - Costa Rica National High Technology Center, Costa Rica
- Bernd Mohr - Jülich Supercomputing Centre, Germany
- David Montoya - Trenz, USA
- Dirk Pleiter - KTH Royal Institute of Technology, Sweden
- Mitsuhsa Sato - Riken, Japan
- Martin Schulz - Technical University of Munich, Germany

Appendix 4 – Attendees & Workshop Photographs

We noted 61 in-person participants plus somewhere between 10 and 19 remote participants. According to the statistics collected by the SC organizers, we had 53 people registered. As noted earlier, the AV crew was not able to get our zoom audio working until mid-way through the Symposium – which affected much of the first half of our program. Still, we Linklings reports that we had 19 Plays, 19 Loads, and 9 Finishes during the week of Supercomputing. We do not have numbers for the rest of November, December, and January.

Last year, SC'21 was mostly virtual and our attendance was 20 in-person participants, 45 remote participants.

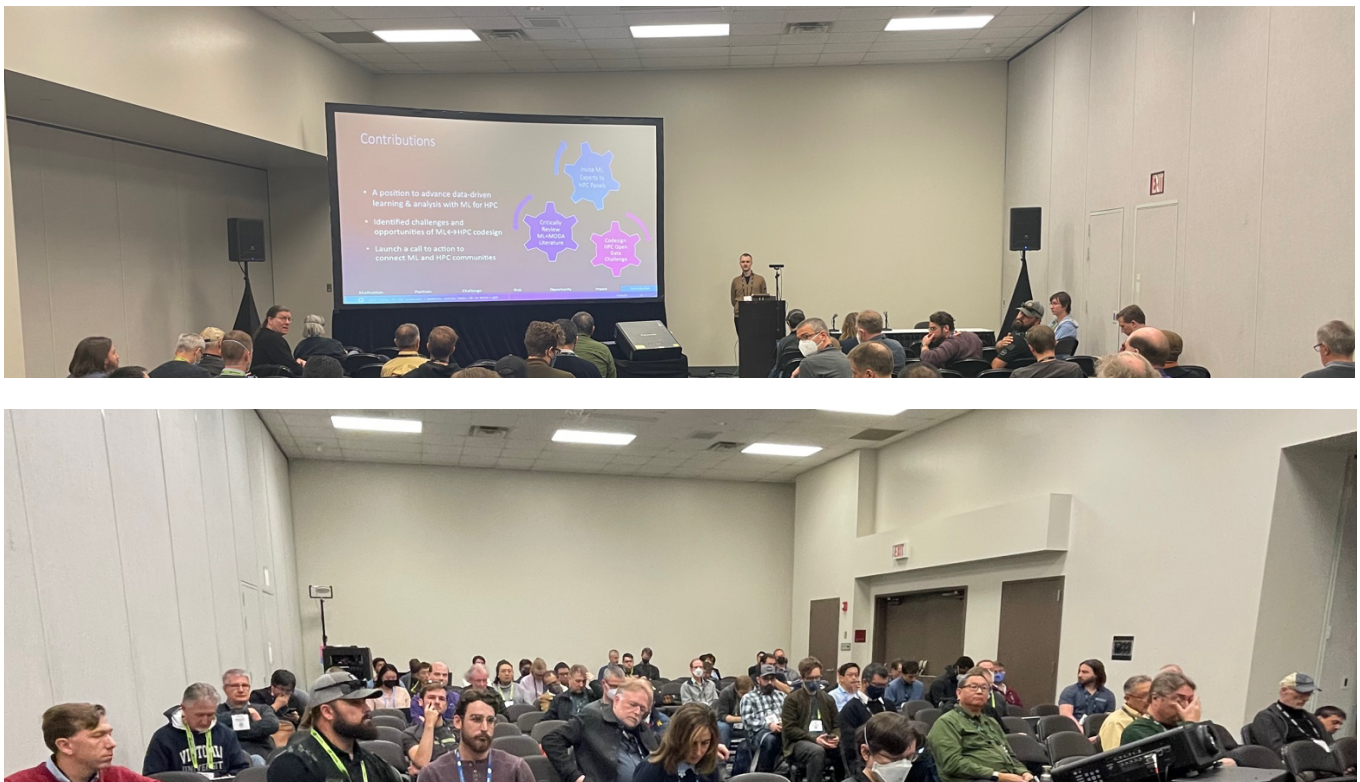


Figure 2 Panoramic views from the 2022 Symposium (top picture is facing meeting front; bottom picture is facing back).



Figure 3 A Question is asked during the Symposium.



Figure 4 Jim Brandt gives a presentation during the Symposium.

Appendix 5 – Position Papers

SQCS 2022 had four position papers:

- William Kramer. “[A Pathway to Achieve the Holy Grail for Efficient Quantitative Co-Design Systems.](#)” Second International Symposium on the Quantitative Codesign of Supercomputers (SQCS 2022), Dallas, TX, Nov. 2022.
- Thomas Jakobsche, Nicolas Lachiche, Florina M. Ciorba. “[Challenges and Opportunities of Machine Learning for Monitoring and Operational Data Analytics in Quantitative Codesign of Supercomputers.](#)” Second International Symposium on the Quantitative Codesign of Supercomputers (SQCS 2022), Dallas, TX, Nov. 2022.
- Brandon Kammerdiener, Michael R. Jantz. “[A Quantitative Approach for Guiding Data Management on Complex Memory Machines.](#)” Second International Symposium on the Quantitative Codesign of Supercomputers (SQCS 2022), Dallas, TX, Nov. 2022.
- Jim Brandt, Ann Gentile. “[AppSysFusion: CoMingling of appropriate data to drive Codesign of Applications, HPC Platforms, and Monitoring, Analysis, and Feedback Infrastructure.](#)” Second International Symposium on the Quantitative Codesign of Supercomputers (SQCS 2022), Dallas, TX, Nov. 2022.

These papers and William Kramer’s presentation are presented below:

Position Paper #1: Thomas Jakobsche

Challenges and Opportunities of Machine Learning for Monitoring and Operational Data Analytics in Quantitative Codesign of Supercomputers

Thomas Jakobsche, PhD Student
University of Basel
thomas.jakobsche@unibas.ch

This work examines the challenges and opportunities of Machine Learning (ML) for Monitoring and Operational Data Analytics (MODA) in the context of Quantitative Codesign of Supercomputers (QCS). MODA is employed to gain insights into the behavior of current High Performance Computing (HPC) systems to improve system efficiency, performance, and reliability (e.g. through optimizing cooling infrastructure, job scheduling, and application parameter tuning [1]).

In this work, we take the position that QCS in general, and MODA in particular, require close exchange with the ML community to realize the full potential of data-driven analysis for the benefit of existing and future HPC systems. This exchange will facilitate identifying the appropriate ML methods to gain insights into current HPC systems and to go beyond expert-based knowledge and rules of thumb. The full potential of ML for QCS is not realized today due to the absence of a set of standard and best practices. To this end, we identify the following **challenges** related to current use of ML for QCS: (a) definition of appropriate operational data to be collected, (b) preparation of data for ML, (c) identification of appropriate ML methods, (d) explainability of ML models, (e) transferability of ML models, (f) FAIR data, privacy and sharing concerns, and (g) data-owners' and machine-learners' perspectives. To address the above challenges, we formulate **opportunities** to bring ML expertise into QCS and facilitate close collaboration: (1) invite ML experts to various thematic panel discussions, (2) review recent advancements in ML, and (3) establish an Open-Data Challenge. HPC administrators, users, and researchers working on improving data-driven operations and quantitative codesign will benefit from the deployment of appropriate ML methods. Advancing ML solutions can leverage the full potential of the vast amounts and types of data being collected on HPC systems. Most MODA solutions in production still involve a human in the loop [2] which prevents the full realization of the vision of autonomous computing systems [7]. ML can enable autonomous responses that are scalable beyond today's manual or ML-assisted capabilities for system optimization.

Challenges

We describe the challenges that need to be addressed to realize the full potential of ML-based data-driven analysis for HPC center-collected data and QCS.

- (a) It is oftentimes challenging for HPC researchers to define what are appropriate data to collect, the individual that a ML model counts/learns on, and the population that such individuals belong to.
- (b) Currently there is no single best way to filter and prepare center-collected data for ML. This is due to the growing data collection capabilities on HPC systems [5][6], resulting in terabytes of high-dimensional, often non-linear, time-series data for each component of the system.
- (c) It is not clear which ML models are suitable for HPC center-collected data, holding back the development of appropriate MODA solutions. There are also no best practices on hyperparameter tuning, performance measuring, suitable training data, and validation in the context of HPC and QCS.

- (d) Absence of explainable ML solutions and missing knowledge behind a specific ML-based decision raise concerns for system operators and prevent them from employing automatic ML-based actions.
- (e) It is an open question whether ML models and insights are transferable outside of the systems they are trained on and if they are useful for the design of next-generation supercomputers.
- (f) HPC researchers face the issues of making data FAIR [8] and resolving privacy and sharing concerns when accessing, analyzing, and publishing data and code.
- (g) Disjoint perspectives and approaches to data analysis: (g.1) Starting with the data, **data owners** ask questions such as: “*What analysis can we do and what problem can we solve?*” and (g.2) Starting with the learning problem, **machine learners** ask questions such as: “*What is the appropriate data and method of analysis?*”. These perspectives need to be reconciled as they play a significant role in the design and development of ML solutions for HPC center-collected data and for QCS.

The **risks** of using ML in QCS without addressing the above challenges are:

- (i) Development of ML-based MODA solutions that only work on curated datasets. These solutions tend to not translate to production systems and real use-cases as their results are often not reproducible.
- (ii) Reluctance to use ML for automated response in production, due to non-explainability.
- (iii) Misleading insights from ML models that result in inappropriate response and/or design decisions.
- (iv) Researchers with a data owner’s perspective will only search for problems and ML methods that fit the data versus finding appropriate data and methods for given problems that ML can help solve.

Opportunities

We identify three opportunities to **bring ML expertise to QCS** and advance data-driven analysis.

Each opportunity is associated with an **objective** and a potential technical approach is also outlined.

- (1) Invite ML experts to directly interact with the QCS communities interested in data-driven analysis of HPC center-collected data via panel discussions held at related annual events such as QCS at SC, HPCMASPA at Cluster, MODA at ISC, PMACS at Euro-Par, and others. Discussion topics can target best practices of collecting and preparing HPC data for ML, identifying appropriate ML models to deploy, and success stories of using explainable ML to autonomously optimize system operations.
- (2) Review recent advancements in ML to reveal appropriate strategies to handle high-dimensional, non-linear, time-series data with a focus on the challenges detailed above. The review can identify areas that are promising for HPC data analysis. Topics to look out for are concept drift, uncertainty quantification, autoencoders, variance thresholding, resampling, dynamic time warping, and time-aggregation. This is an opportunity for all active researchers in the ML and QCS communities.
- (3) Initiate an “Open-Data Challenge” (ODC) to facilitate the development of ML for QCS. We envision an anomaly detection challenge based on representative datasets from different centers. The datasets can be generated through executions of representative (proxy-)applications and the performance anomalies produced by the HPC Performance Anomaly Suite [4]. The design of the challenge should include ML experts in order to correctly define the format of the challenge, the desired insights about ML for HPC, and how the challenge should be formulated to get these insights. This challenge is also a way to address the above mentioned concerns regarding data sharing and transferable ML models. ODC can spawn reproducible work and reveal whether different approaches

and ML models are necessary to address similar use-cases on different HPC systems. ODC can also include a call to explore portable ML models across multiple systems [3].

Conclusion

This position paper is an initial effort and a call to the HPC and ML communities to connect and collaboratively build and deploy best and standard practices of using ML for QCS. The close exchange between the QCS and ML communities will reconcile the data owners' and machine learners' perspectives and lead to the development of appropriate ML-based solutions that will benefit and improve the operations of existing and the design of future HPC systems. A first immediate action is to invite ML experts to upcoming events in QCS and MODA. A subsequent immediate action is to organize, launch, and evaluate the ODC.

Table 1: Complete author list and institutional affiliations

Position Paper Author List			Institution
Jakobsche	Thomas	PhD Student	University of Basel, Switzerland
Lachiche	Nicolas	Professor	University of Strasbourg, France
Ciorba	Florina M.	Professor	University of Basel, Switzerland

Optional References Here

- [1] Netti, A., Shin, W., Ott, M., Wilde, T., & Bates, N. "A conceptual framework for hpc operational dataanalytics." 2021 IEEE International Conference on Cluster Computing (CLUSTER). IEEE, 2021.
- [2] Ott, M., Shin, W., Bourassa, N., Wilde, T., Ceballos, S., Romanus, M., & Bates, N. "Globalexperiences with HPC operational data measurement, collection and analysis." 2020 IEEE International Conference on Cluster Computing (CLUSTER). IEEE, 2020.
- [3] Netti, A., Ott, M., Guillen, C., Tafani, D., & Schulz, M. "Operational Data Analytics in practice: Experiences from design to deployment in production HPC environments." *Parallel Computing* 113 (2022): 102950.
- [4] Ates, E., Zhang, Y., Aksar, B., Brandt, J., Leung, V. J., Egele, M., & Coskun, A. K. "HPAS: An HPCperformance anomaly suite for reproducing performance variations." *Proceedings of the 48th International Conference on Parallel Processing*. 2019.
- [5] Agelastos, A., Allan, B., Brandt, J., Cassella, P., Enos, J., Fullop, J., ... & Tucker, T. "The lightweightdistributed metric service: a scalable infrastructure for continuous monitoring of large scale computing systems and applications." *SC'14: Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis*. IEEE, 2014.
- [6] Netti, A., Müller, M., Auweter, A., Guillen, C., Ott, M., Tafani, D., & Schulz, M. "From facility toapplication sensor data: modular, continuous and holistic monitoring with DCDB."

Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis. 2019.

- [7] Horn, P. "Autonomic computing: IBM's perspective on the state of information technology." (2001).
- [8] Wilkinson, M. D., Dumontier, M., Aalbersberg, I. J., Appleton, G., Axton, M., Baak, A., ... & Mons, B. "The FAIR Guiding Principles for scientific data management and stewardship." Scientific data 3.1 (2016): 1-9

Position Paper #2: Brandon Kammerdiener

A Quantitative Approach for Guiding Data Management on Complex Memory Machines

Brandon Kammerdiener and Michael R. Jantz

University of Tennessee, Knoxville

(423) 863-0807, [bkammerd,mrjantz]@vols.utk.edu

Introduction

Proliferation of real-time and AI-driven decision making continues to fuel the need for ever faster access to larger sets of data in memory. At the same time, increasing demands for high-density sharing are leading to more supercomputing configurations with large amounts of memory attached to each node and connected through efficient networking resources. New media technologies, such as high bandwidth memories, phase change memories, spin-torque transfer RAMs, and many others, and new processing-memory interconnect options, including the Compute Express Link (CXL), are bringing rich opportunities for addressing the diverse needs of applications under various cost, performance, and power constraints. In response to these trends, many supercomputing systems now include a heterogeneous mix of memory devices and organizations, which can enable the combined benefits of their unique capabilities and support such diverse and multi-tenant workloads in modern computing centers.

Despite their potential benefits, heterogeneous memory architectures present new challenges for data management. Computing systems have traditionally viewed memory as a single homogeneous address space, sometimes divided into different non-uniform memory access (NUMA) domains but consisting entirely of the same storage medium (i.e., DDR* SDRAM). To utilize heterogeneous resources efficiently, alternative strategies are needed to match data to the appropriate technology in consideration of hardware capabilities, application usage, and in some cases, NUMA domain.

Spurred by this problem, the architecture and systems communities have proposed a range of hardware and software techniques to manage data efficiently on heterogeneous memory systems. For example, some systems choose to utilize high performance memories as a large, hardware-managed cache. Such configurations are not only inflexible, but they are also often inefficient because they require storage for very long tag values and additional bandwidth to distinguish cache blocks at DRAM scale. Alternatively, many systems employ a software-based approach, where the operating system (OS), or the OS in collaboration with the applications, assign and move data into different memory and storage devices. Many modern platforms also provide system-level interfaces and custom allocators that allow the applications themselves to control the placement of data across the memory hierarchy. While software-based controls increase flexibility, they are still limited because they either proceed with little or no knowledge of how the applications intend to use memory resources or they require developers with detailed knowledge of complex memory resources and the capability and resources to update existing applications. These constraints are particularly problematic for scientific and high-performance computing (HPC) applications due to the frequency and scale of data usage as well as the need for performance portability in the face of a rapidly changing architectural landscape.

Our Approach

Building upon our own prior work in heterogeneous memory management [1 – 6], we propose to address these challenges through a quantitative approach that leverages detailed profiling and analysis of application behavior to steer data management of complex memory platforms. Our proposed approach is fully automatic and enables guided data tiering of individual program data objects, without requiring application modification or even recompilation. Figure 1 depicts the main components of our approach, which primarily consists of two new pieces of user-level software:

- 1) Custom runtime facilities, which enable applications to invoke the monitoring and management capabilities of our framework without updates to program source code, and
- 2) A systemwide monitoring and data management process, which runs alongside the applications and conducts object tiering through a series of complementary activities, including (1) monitoring and structuring profiles of object allocation and usage, (2) automated heuristics to prioritize objects for placement in fast memory, and (3) mechanisms to enforce tier recommendations when a particular event occurs.

By providing automated monitoring and analysis along with customizable controls, this approach will enable applications to define custom quantitative strategies for guiding data tiering on heterogeneous memory machines. Moreover, it provides these capabilities by leveraging standard tools and system interfaces and can therefore be deployed in real and diverse supercomputing environments, with minimal effort. Thus, we believe that this approach has significant potential to address the challenges presented by complex memory hierarchies and will enable more efficient execution for a wide range of supercomputing applications.

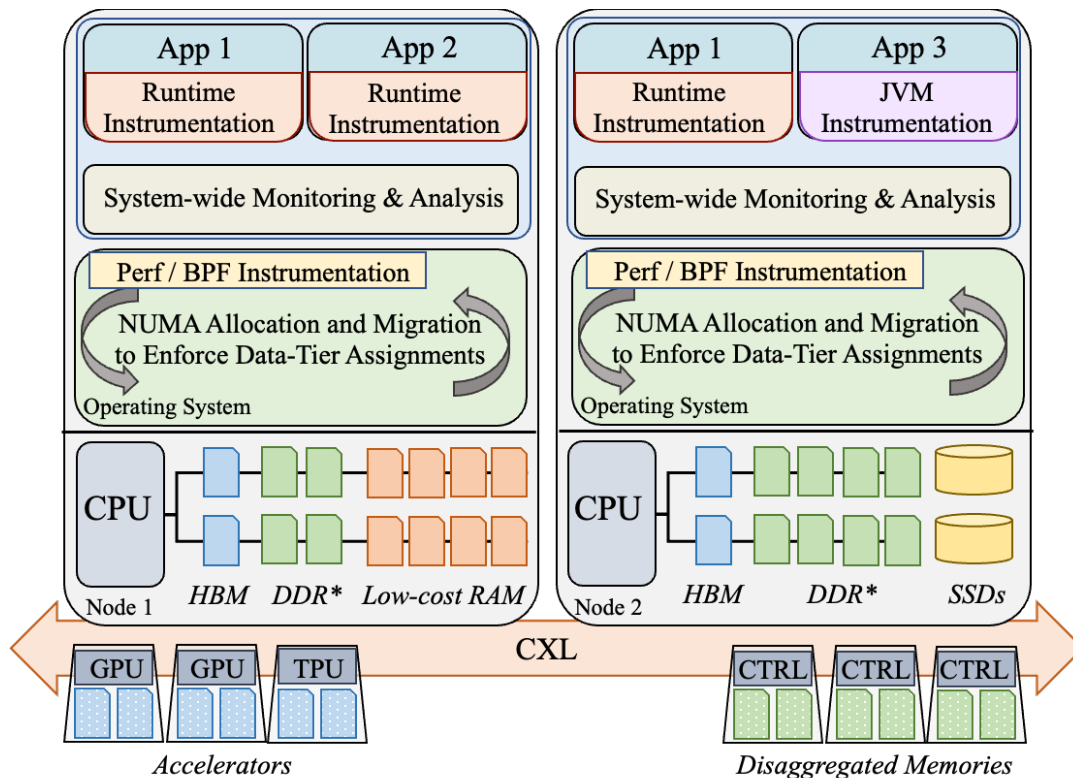


Figure 1. Design overview of our approach

Table 1: Complete author list and institutional affiliations

Position Paper Author List			Institution
Last Name	First Name	Title	Institution Name
Kammerdiener	Brandon	Graduate Research Assistant	University of Tennessee
Jantz	Michael R.	Associate Professor	University of Tennessee

Optional References Here

1. Effler, T. C., Howard, A. P., Zhou, T., Jantz, M. R., Doshi, K. A., and Kulkarni., P. A. On automated feedback-driven data placement in hybrid memories. In LNCS International Conference on Architecture of Computing Systems (April 2018).
2. Effler, T. C., Jantz, M. R., and Jones, T. Performance potential of mixed data management modes for heterogeneous memory systems. In 2020 IEEE/ACM Workshop on Memory Centric High Performance Computing (MCHPC) (Washington, DC, USA, 2020), pp. 10–16.
3. Effler, T. C., Kammerdiener, B., Jantz, M. R., Sengupta, S., Kulkarni, P. A., Doshi, K. A., and Jones, T. Evaluating the effectiveness of program data features for guiding memory management. In The International Symposium on Memory Systems (2019), MEMSYS '19, p. 383– 395.
4. Olson, M., Kammerdiener, B., Jantz, M., Doshi, K., and Jones, T. Portable application guidance for complex memory systems. In The International Symposium on Memory Systems (09 2019), MEMSYS '19, ACM, pp. 156–166.
5. Olson, M. B., Kammerdiener, B., Doshi, K. A., Jones, T., and Jantz, M. R. Online application guidance for heterogeneous memory systems. To appear in the ACM Transactions on Architecture and Code Optimization (May 2022).
6. Olson, M. B., Zhou, T., Jantz, M. R., Doshi, K. A., Lopez, M. G., and Hernandez, O. R. Membrain: Automated application guidance for hybrid memory systems. In IEEE Intl. Conf. on Networking, Arch. and Storage, NAS, October 11-14, 2018. Best Paper Award. (2018), pp. 1–10.

Position Paper #3: Jim Brandt

AppSysFusion: CoMingling of appropriate data to drive Codesign of Applications, HPC Platforms, and Monitoring, Analysis, and Feedback Infrastructure

James Brandt, Technical Staff

Sandia National Laboratories

(925) 519-1999, brandt@sandia.gov

Key Points

- **Need to move from considering individual application performance to considering the efficiency of the multi-dimensional HPC ecosystem (i.e., optimize utilization of resources and performance while honoring constraints such as power and priority)**
- **The dynamic and complex nature of HPC workloads requires continuous orchestration of overall HPC ecosystem which relies on continuous insight into all dimensions**
- **We must re-imagine HPC resources as autonomous peer components that can negotiate among themselves and with applications to optimize global efficiency**

Background

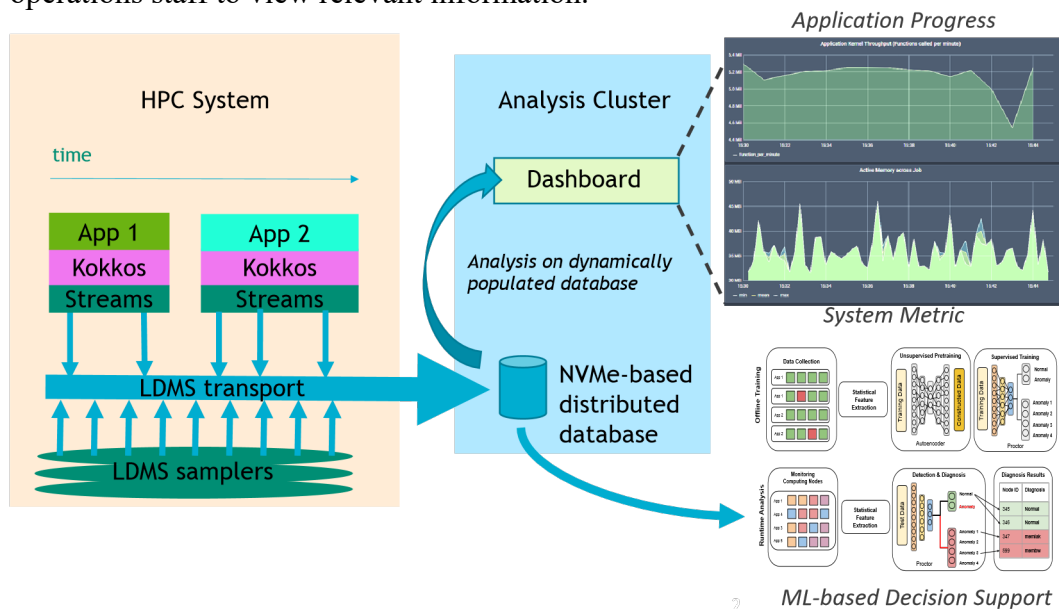
The goal of building HPC systems is to enable execution of large-scale user application workflows in an efficient and performant manner. Performance here is multi-dimensional and includes not just a particular application's time-to-solution, but the aggregate throughput of all applications submitted (workload) and energy spent. The aggregate HPC system power draw must always remain within a contract envelope. Note that workflow in this context refers to an individual user's end-to-end pipeline of executions (some perhaps executed in parallel) and not necessarily a single application run (e.g., iterating a number of times through parameter selection, simulation, results analysis, and visualization). A workload is defined as the combined set of jobs concurrently being executed on an HPC system including flushing of buffers to stable storage and pre-staging of new data to buffers.

Given this efficiency and performance goal, associated constraints, and possible relative prioritization of individual applications and workflows, data-driven scheduling, resource allocation, and feedback will be key to achieving it. Appropriate data from both workflow and system components needs to be acquired and processed in order to make informed resource scheduling and allocation decisions. Additionally, there must be capabilities built into applications and system software to enable them to accept and respond to feedback. Processing of acquired data must be performed on actionable timescales to be of benefit to feedback-enabled applications and system software.

Temporally coherent data collected from across system and workload components is necessary for gaining understanding of how workload components (i.e., individual applications of disparate workflows) contribute to resource load, affect resource performance, and respond to the effects of resource performance changes. As resources become more heterogeneous gaining such understanding (e.g., is it better to execute a particular phase of a workflow on CPU, GPU, FPGA, or? What storage should be used? When, if at all, to checkpoint?) becomes more difficult due to the many options and the need to port applications to these different technologies to obtain technology relevant data. Additionally, understanding how resources respond to oversubscription and how that affects individual execution phases of individual applications, given aggregate resource demand is difficult even in today's workload environment. This environment is only becoming more complex over time as new HPC applications emerge to share the already crowded HPC ecosystem.

Current Approach

To address these challenges, Sandia has launched its AppSysFusion project (shown diagrammatically below). The project comprises functional elements to: 1) collect temporally coherent data with associated absolute timestamp information from all workload components (e.g., application progress/performance measures) and monitorable system elements (e.g., synchronously collect compute node, network, and storage parameters), 2) aggregate collected data to a common distributed data store, 3) process data (e.g., statistical analysis, Machine Learning (ML), generate feedback data) as it arrives to the data store, and 4) provide a visualization portal for users and operations staff to view relevant information.



While AppSysFusion enables collection of application data (via Kokkos depicted above), a consistent method of labeling {application, decomposition, technology, parameter, xxx} combinations is needed to construct and identify models that can be used to allocate “good-fit” resources, ensure resources aren’t overloaded, enable automated detection of inefficient behavior, and provide automated feedback to guide more performant/efficient execution/resource utilization. Further, a well-defined information pipeline to enable users and system administrators to evaluate workload, workflow, and resource interactions and utilizations is needed to provide appropriate feedback to application developers and system architects about needed technology directions, instrumentation, and feedback hooks.

Work in Progress

Fundamental changes in resource accounting and management are also required to create the symbiotic relationships and communication paths among workload/workflow components, and HPC resources that can naturally drive performance, efficiency, and throughput toward optimality. Identifying gaps in properties of information and mechanisms for sharing will be needed for co-design of next-generation platforms to further augment hardware and software related performance gains. Challenges to taking this approach are development and testing of: 1) software and hardware to enable scalable distributed decision-making at HPC resources (e.g., spare cycles or dedicated processors on compute nodes, service nodes, and storage devices). Not a dedicated monitoring

system) without negatively impacting resources executing user workflow performance, 2) a communication protocol that is lightweight and can accommodate required communication and latency bounds, and 3) system software for birth-to-death management of workflows including addition of self-aware workflow components that can interact in a run time customer-provider relationship with distributed intelligent resource components. Our current approach is to create a container-based emulator in which we can develop and explore interaction of the above-described software components, in the context of a large-scale emulated system, and utilize labeled data traces from production systems as ground-truth for workflow component behavior and resource requirements.

Brandt, AppSysFusion:CoMingling and CoDesign

2

The position paper must include the following Table, which will not count toward the two-page limit. No abbreviations or acronyms should be included in the table.

Table 1: Complete author list and institutional affiliations

Position Paper Author List			Institution
Last Name	First Name	Title	Institution Name
Brandt	James	Technical Staff	Sandia National Laboratories
Gentile	Ann	HPC Manager	Sandia National Laboratories

References are optional but recommended

Optional References Here

Invited Position Presentation: William Kramer



1

The Good News

- Many supercomputing organizations perform system evaluation and analysis with continuously collected system-wide data collection
- combined with vendor independent, community SW tools to do real-time system management, application performance improvements all the way to long term Quantitative Co-Design. When merged with models and/or kernels of the next generation of applications and methods, we may be able to rapidly and fully evaluate many configuration and systems to optimize the next generation technologies.

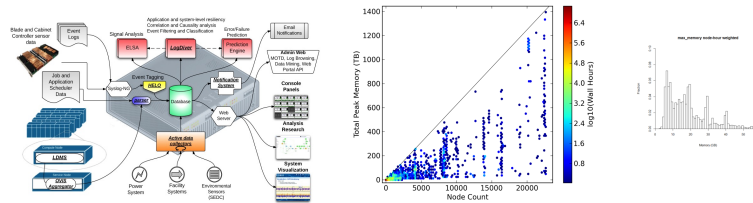


NCSA | NATIONAL CENTER FOR SUPERCOMPUTING APPLICATIONS

2

The Good News

- Many supercomputing organizations perform system evaluation and analysis with continuously collected system-wide data collection



NCSA | NATIONAL CENTER FOR SUPERCOMPUTING APPLICATIONS

3

Examples of what can be done

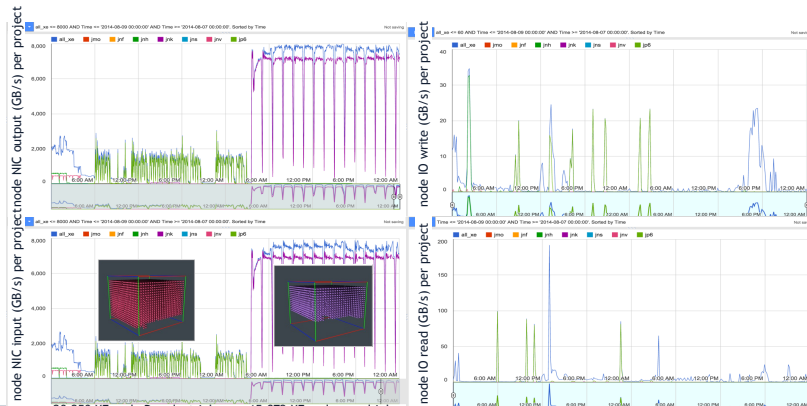
- LDMS deployed at scale (> 11M data points per minute) on Petascale Systems without introducing Jitter
 - Lightweight Distributed Metric Service: A Scalable Infrastructure for Continuous Monitoring of Large Scale Computing Systems and Applications, A. Agelastos, B. Allan, J. Brandt, P. Cassella, J. Enos, J. Fullop, A. Gentile, S. Monk, N. Naksinehaboon, J. Ogden, M. Rajan, M. Showerman, J. Stevenson, N. Taerat, and T. Tucker
 - IEEE/ACM Int'l Conf. for High Performance Storage, Networking, and Analysis (SC14) New Orleans, LA, Nov 2014.
- Example Insights from ISC and Logdiver
 - 99.4% of failures limited to a single blade;
 - Software errors propagate 20 times more often than hardware failures;
 - DDR5 ECC is 100x more prone to uncorrected errors than DDR3 with x8 Chipkill;
 - software accounts for 53% of repair hours;
 - hardware failure rates decline over time but software errors do not;
 - 74% of system wide outages are due to software;
 - 50% of the software failures occur are during failover;
 - filesystem and interconnect are prime contributors;
 - failure of failover causes a significant number of system wide outages;
 - application failure increases with increasing duration of failover time.
 - Catello Di Martino, Zbigniew Kalbarczyk, William Kramer, Ravishankar Iyer, "Measuring and understanding extreme-scale resilience: A field study of 5,000,000 HPC application runs," 45th IEEE/IFIP International Conference on Dependable Systems and Networks, DSN 2015, pp. 25–36, 22–25 June 2015. URL: <http://ieeexplore.ieee.org/stamp/stamp.jsp?tp=&arnumber=7266835>
 - Di Martino, Catello, F. Baccanico, W. Kramer, J. Fullop, J. Z Kalbarczyk, and R Iyer, Lessons Learned From the Analysis of System Failures at Petascale: The Case of Blue Waters, The 44th Annual IEEE/IFIP International Conference on Dependable Systems and Networks (DSN 2014)), June 23-26 2014



NCSA | NATIONAL CENTER FOR SUPERCOMPUTING APPLICATIONS

4

Examples of What We Do Today



NCSA | NATIONAL CENTER FOR SUPERCOMPUTING APPLICATIONS

5

Example – Understanding LS-DYNA Performance 200M DOF and 2,048 MPI ranks

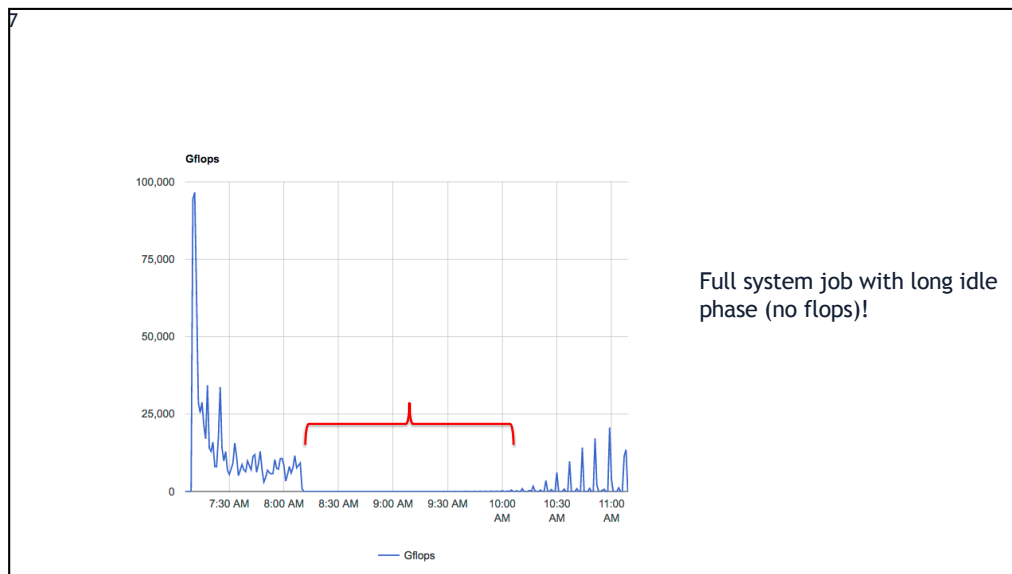
The spike is
MPI_ALLREDUCE
The redistribution
spike is
asynchronous
sends/receives
(we also have an
MPI_ALLTOALLV
variant).
The factorization
"ramp up" is
largely MPI_BCAST.



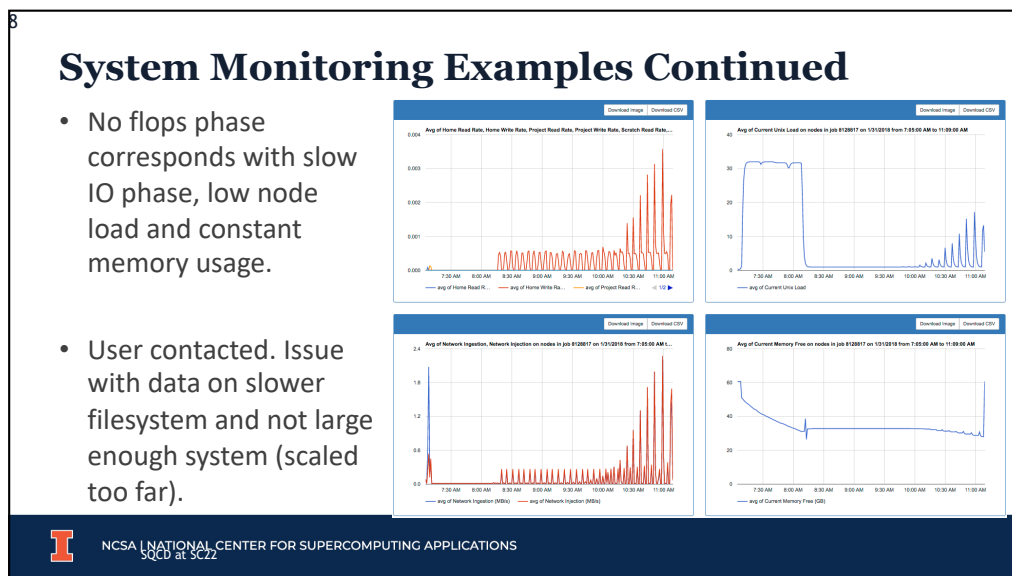
NCSA | NATIONAL CENTER FOR SUPERCOMPUTING APPLICATIONS

SQCD at SC22

6



7



8

So, Why Don't We Do This Every Time for Everything?

- While the raw data collection is reasonably portable (e.g. LDMS, Lustre Stats,...) there is a wide variation in implementations
 - Site specific configurations
 - Site specific storage
 - Dashboard vs analysis usage
 - Site/system specific jobs
- Lack of shared data
- Us of ML/AI limited and hard

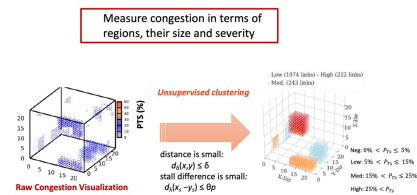


NCSA | NATIONAL CENTER FOR SUPERCOMPUTING APPLICATIONS

9

Example – AI integrated analysis

- HSN congestion is the biggest contributor to app performance variation
- Continuous presence of high congestion regions
- Long lived congestion (may persist for >23 hours)
- Default congestion mitigation mechanism have limited efficacy
- Only 8 % (261 of 3390 cases) of high congestion cases found using our framework were detected and acted by default congestion mitigation algorithm
- In ~30% of the cases the default congestion mitigation algorithm was unable to alleviate congestion
- Congestion patterns and their tracking enables identification of culprits behind congestion critical to system and application performance improvements
 - E.g., intra-app congestion can be fixed by changing allocation and mapping strategies



Measuring Congestion in High-Performance Datacenter Interconnects

Surabhi Iyer, Archit Patke, Jim Brandt, Ann Centile, Benjamin Lim, Mike Showerman, Greg Bauer, Larry Kaplan, Zbigniew Kalbarczyk, William Kramer, Ravi Iyer
In Proceedings of the 17th USENIX Symposium on Networked Systems Design and Implementation (NSDI 20)

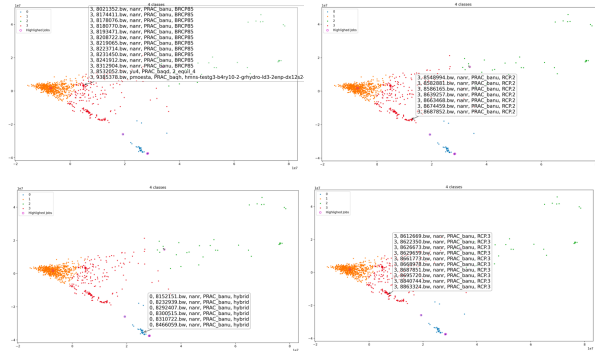


NCSA | NATIONAL CENTER FOR SUPERCOMPUTING APPLICATIONS

10

Example: Latent Space projected scatter plot of Autoencoding-RNN

- Two major challenges of system monitoring data this scale and complexity.
 - Terabytes of data have no labels,
 - Dimensionality of each sample was on the order of billions.
- Unsupervised training is required
- 1600 cores for many hours



From unpublished work by Arron Saxton and Albert Bode under funding from DOE Tri-labs



NCSA | NATIONAL CENTER FOR SUPERCOMPUTING APPLICATIONS

11

How Do We Progress to Holistic, Quantitative Codesign and System Management?

- Much more sharing of real data
 - By the end of the year Blue Waters will release 9 years of system and performance data for anyone to use.
- Better and more accurate ways to anonymize collected data
 - Should be open source for all to use
 - Should work for National Lab resources
- Labelled data repositories (aka Resnet) for monitoring data
- Data Collection is not the issue, automated understanding is
 - Piecemeal solutions will not get there in time



NCSA | NATIONAL CENTER FOR SUPERCOMPUTING APPLICATIONS

12