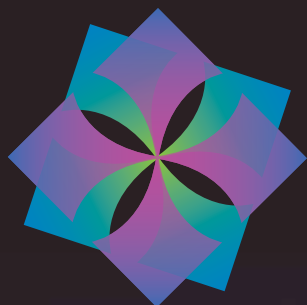


Ernest N. Morial Convention Center • New Orleans, Louisiana



SC14

New Orleans, LA | **hpc**
matters.

Exhibition Dates:
November 17-20, 2014
Conference Dates:
November 16-21, 2014

Conference Program

The International Conference for High Performance Computing, Networking, Storage and Analysis

Sponsors:



Association for
Computing Machinery





SC14
New Orleans, LA | **hpc**
matters.

Table of Contents

3 Welcome from the Chair

4 SC14 Mobile App

5 General Information

9 SCinet Contributors

11 Registration Pass Access

13 Maps

16 Daily Schedule

26 Award Talks/Award Presentations

31 Birds of a Feather

50 Doctoral Showcase

56 Exhibitor Forum

**67 HPC Impact Showcase/
Emerging Technologies**

82 HPC Interconnections

92 Keynote/Invited Talks

106 Papers

128 Posters

152 Tutorials

164 Visualization and Data Analytics

168 Workshops

178 SC15 Call for Participation

Welcome to SC14



HPC helps solve some of the world's most complex problems. Innovations from our community have far-reaching impact in every area of science—from the discovery of new drugs to precisely predicting the next superstorm—even investment

banking! For more than two decades, the SC Conference has been the place to build and share the innovations that are making these life-changing discoveries possible.

This year, SC returns to New Orleans with new ideas and a fresh approach to HPC. Spotlighting the most original and fascinating scientific and technical applications from around the world, SC14 will once again bring together the HPC community – an unprecedented array of scientists, engineers, researchers, educators, students, programmers, system administrators, and developers – for an exceptional program of technical papers, tutorials, timely research posters, panels and workshops, and Birds-of-a-Feather sessions.

Long-time HPC community members know SC is the hub of our community. While there are many, many excellent technical conferences in HPC, the SC Conference series, sponsored each year by the Association for Computing Machinery and the IEEE Computer Society, is the largest and most diverse conference in our community. Recent conferences have brought together more than 10,000 attendees during conference week from all over the world, truly making it the international conference for high performance computing, networking, storage, and analysis.

This year we're expecting more than 350 exhibitors spread over 151,000 square feet of exhibit space. With 78 meeting rooms dedicated to the Technical Program, Exhibitor Forum, Tutorials, Workshops, and HPC Interconnections, and several large rooms for Technical Program events, there'll be plenty of space for all of our events, yet laid out in a way that makes it quick and easy to find your way around.

SC is fundamentally a technical conference, and anyone who has spent time on the show floor knows the SC Exhibits program provides a unique opportunity to interact with the future of HPC. Far from being just a simple industry exhibition, our research and industry booths showcase recent developments in our field, with a rich combination of research labs, universities, and other organizations and vendors of all types of software, hardware, and services for HPC.

Another defining characteristic of SC is SCinet, the conference network that, for the week of the conference, is one of the largest and most advanced networks in the world. While many attendees experience SCinet directly as wireless networking provided throughout the convention center, SCinet also provides a once-a-year opportunity for research and engineering groups to work with some of the most modern networking equipment and very high-bandwidth wide-area networking. This year, SCinet will provide nearly 1 Terabit/second of networking bandwidth! During the conference please be sure to stop by the SCinet booth for more information on the networking infrastructure.

For over a quarter of a century the SC Conference has served as an inflection point for the entire HPC community. New Orleans is a city of historical significance, of great pride and of reinvention in the face of great challenges. I can't think of a better place to celebrate SC's rich past and our community's ever-evolving future.

Welcome to New Orleans and SC14, where HPC Matters!



Trish Damkroger
SC14 General Chair



SC14 Mobile App*

Want an easy way to keep up with all that SC14 offers? Install the new SC14 Mobile App on your Apple or Android device and customize your experience!

Search for Supercomputing Conference Series in your app store or scan this QR Code on your mobile device to get started today.



The SC14 app includes standard components from years past, but with new features to improve your attendee experience. New features include:

- Searchable event schedule for easy-to-find sessions and events
- Attendee personalization, scheduling and notes within the app
- Local guide to help you get around
- SC14 updates to keep you informed and involved
- Exhibitor directory with map
- Ability for attendees to easily share their SC experience through social media
- QR Code reader to quickly find SC14 information
- Integration with ACM Digital Library to easily find session speakers, papers, and abstracts during the conference.

To use the app simply log in with your email address and your SC14 Confirmation # as your password. Logging in will allow your calendar, notes, and other in-app data to sync across multiple devices.

You can even access your data online at: <http://visitors.genie-connect.com/sc14> to modify your schedule, export your schedule to your Outlook, iCal or Google calendar, and create a complete report of your SC14 activities after the conference.

We hope your SC14 app experience is as great as the conference itself. If you have any questions, comments, or problems, please email us at sc14app@info.supercomputing.org.

*This app is developed through an agreement with SC14 and GenieConnect.

General Information

General Information

In this section, you'll find information on registration, exhibit hours, conference store hours, description and locations of all conference social events, information booths and their locations, as well as convention center facilities and services.

General Information

Registration and Conference Store

The registration area and conference store are located in the Lobby H of the New Orleans Convention Center.

Hours:

Saturday, November 15	1pm – 6pm
Sunday, November 16	7am – 6pm
Monday, November 17	7am – 9pm
Tuesday, November 18	7:30am – 6pm
Wednesday, November 19	7:30am – 6pm
Thursday, November 20	7:30am – 5pm
Friday, November 21	8am – 11am

Registration Pass Access

See pages 11 and 12 for access grids.

Exhibit Hall Hours

Tuesday, November 18	10am-6pm
Wednesday, November 19	10am-6pm
Thursday, November 20	10am-3pm

SC14 Information Booth

Need up-to-the-minute information about what’s happening at the conference. Need to know where to find your next session? What restaurants are close by? Where to get a document printed? These questions and more can be answered by a quick stop at one of the SC Information booths. There are two booth locations for your convenience: the main booth is located in Lobby H; the second booth is located in Lobby F/G.

Information Booth hours are as follows (note that the times in parentheses indicate hours for the secondary booth; closed if no time is indicated):

Saturday	1pm-6pm
Sunday	8am-6pm (8am-5pm)
Monday	8am-9pm (8am-5pm)
Tuesday	7:30am-6pm (7:30am-4pm)
Wednesday	8am-6pm (8am-4pm)
Thursday	8am-6pm
Friday	8:30am-12pm

SC15 Preview Booth

Members of next year’s SC committee will be available in the SC15 preview booth (located in Lobby I, across from the conference store) to offer information and discuss next year’s SC conference in Austin, Texas. Stop by for a copy of next year’s *Call for Participation* and pick up some free gifts!

The booth will be open during the following hours:

Tuesday, Nov. 17	10am-4pm
Wednesday, Nov. 18	10am-4pm
Thursday, Nov. 29	10am-4pm

Social Events

Exhibitors’ Party

Sunday, November 16
6pm-9pm

Audubon Aquarium of the Americas (bus departs from Lobby I)

The exhibitor reception is SC14’s way of thanking exhibitors for their participation and support of the conference. The event will feature local band *Creole String Beans*, access to all the displays and exhibits at the aquarium, and a showcase of local food and drinks throughout the facility. Transportation will be provided from Lobby I of the convention center starting at 5:45pm; last bus returns at 9pm.

An Exhibitor badge and party ticket are required to attend this event.

Exhibits Gala Opening Reception

Monday, November 17
7pm-9pm

SC14 will host its annual Grand Opening Gala in the Exhibit Hall. This will be your first opportunity to see the latest high performance computing, networking, storage, analysis, and research products, services, and innovations. This event is open to all Technical Program, Exhibitors, and HPC Interconnections Program (Broader Engagement, HPC for Undergraduates and Student Cluster Competition) registrants.

Posters Reception

Tuesday, November 18
5:15pm-7pm

The Posters Reception will be held in the New Orleans Theater Lobby. The reception is an opportunity for attendees to interact with poster presenters and research and ACM Student Research Competition (SRC) posters. The reception is open to all attendees with Technical Program registration. Complimentary refreshments and appetizers are available until 7pm.

Technical Program Conference Reception

Thursday, November 20

6pm-9pm

The Sugar Mill

The Technical Program Conference reception is intended to give the Technical Program attendees a chance to network and collaborate in a relaxed setting. Join us for a taste of New Orleans street food, local beverages, and entertainment. A Technical Program badge, event ticket, and government-issued photo ID are required to attend this event. Attendees are required to wear technical program badges throughout the reception, and badges may be checked during the event. The Sugar Mill is within walking distance of the convention center. Shuttle transportation between the hotels and convention center will run until 9pm from Lobby I during the event.

ATMs

ATMs are located outside Halls E and I.

Business Center (UPS Store)

The Business Center is located in Lobby F.

City and Dining Information

City and dining information are provided by “On the Town” concierge located next to the main Information Booth (Lobby H).

Coat & Bag Check

The coat and bag check station is located in Lobby H. The hours are as follows:

Saturday, November 15	7am-5:30pm
Sunday, November 17	7am-6pm
Monday, November 18	7:30am -9:30pm
Tuesday, November 19	7:30am-7:30pm
Wednesday, November 20	7:30am-7:30pm
Thursday, November 21	7:30pm-7:30pm
Friday, November 22	7:30pm-3:30pm

Emergency Contact

For an onsite emergency, please dial 3040 on any house phone.

Family Day

Family Day is Wednesday, November 20, 4pm-6pm. Adults and children 12 and over are permitted on the floor during these hours when accompanied by a registered conference attendee.

First-Aid Center

There are two first aid offices in the convention center. One is on the east side, located next to Room 150A; the other is on the west side lobby outside Hall H.

Lost Badge

There is a \$40 processing fee to replace a lost badge.

Lost & Found

Lost and found is located in Room 258.

Prayer & Meditation Room

The prayer and meditation room is located in Room 254 of the convention center.

Restrooms

Restrooms are located conveniently throughout the convention center. Locations are colored pink on the convention center brochure, which can be picked up at the SC14 Information Booths.

Wheelchair/Scooter Rental

Wheelchairs and medical mobility scooters are available for rent at the Business Center.

SCinet

SCinet strives to make the best wireless network resources available for you throughout the conference. Each year we are making improvements to the service, working in partnership with the city of New Orleans and the convention center IT staff. All of SC's networking resources—including nearly a terabit of wired bandwidth—are provided through SCinet. This year, SCinet will deploy IEEE 802.11a, 802.11g and 802.11n wireless networks that will enable basic access to the Internet at no charge to attendees. The wireless network will extend to meeting rooms, exhibit halls, and all common areas of the center occupied by SC14.

SC attendees are among the most connected conference-goers around the world and put even the most sophisticated wireless networks to the test. We apologize in advance if you experience any connection difficulties in certain areas due to reduced signal strength, network congestion, or other issue. SCinet actively monitors the health of the wireless networks and will make every effort to provide stable and reliable services. If you have any questions during the conference, please contact us at scinet@info.supercomputing.org.

Additional Wireless Policy Details

In order to provide the most robust wireless service possible, SCinet has some special policies in place. SCinet must control the entire 2.4GHz and 5.2GHz ISM bands (2.412GHz to 2.462GHz and 5.15GHz to 5.35GHz) within the convention center. This has important implications for exhibitors and attendees:

- Exhibitors and attendees may not operate their own IEEE 802.11 (a,b,g,n or other standard) wireless Ethernet access points anywhere within the convention center, including within their own booth.
- Wireless clients may not operate in ad-hoc or peer-to-peer mode due to the potential for interference with other wireless clients.
- Exhibitors and attendees may not operate 2.4GHz or 5.2GHz cordless phones or microphones, wireless video or security cameras, or any other equipment transmitting in the 2.4GHz or 5.2GHz spectrum.

SCinet wants everyone to have a successful, pleasant experience at SC14. This should include the ability to sit down with your wireless-equipped laptop or PDA to check e-mail or surf the web from anywhere in the wireless coverage areas. Please help us achieve this goal by not operating equipment that will interfere with other users.

The SC14 Conference reserves the right to deny or remove access to any system in violation of the SCinet acceptable usage policy and to anyone who uses multicast applications or harms the network in any way, intended or unintended, via computer virus, excessive bandwidth consumption or similar misuse.

SCinet Contributors



SCinet, the world's fastest network, would not be possible without the generosity of our vendors and service providers.

Thank you to the SCinet Contributors of SC14!

P L A T I N U M



SCinet Contributors



SCinet, the world's fastest network, would not be possible without the generosity of our vendors and service providers.

Thank you to the SCinet Contributors of SC14!

G O L D



S I L V E R



B R O N Z E



Registration Pass Access

Registration Pass Access - Technical Program

Each registration category provides access to a different set of conference activities, as summarized below.

Type of Event	Tutorials	Technical Program	Technical Program + Workshops	Workshop Only
Awards (Thursday)		*	*	*
Birds-of-a-Feather	*	*	*	*
Broader Engagement (Sun/Mon)			*	*
Broader Engagement (Tue-Thu)		*	*	
Conference Reception (Thursday)		*	*	
Exhibit Floor		*	*	
Exhibitor Forum		*	*	
Exhibits Gala Opening (Monday)		*	*	
HPC Matters Plenary (NO BADGE REQUIRED/OPEN TO PUBLIC)	*	*	*	*
Invited Talks (Non-Plenary)		*	*	
Invited Talks (Plenary)		*	*	
Keynote (Tuesday)	*	*	*	*
Panels (Tue-Thur)		*	*	
Panels (Friday Only)		*	*	
Papers		*	*	
Posters		*	*	
Poster Reception (Tuesday)		*	*	
Tutorial Lunch (Sun/Mon ONLY)	*			
Tutorial Sessions	*			
Student Cluster Competition		*	*	
Workshops			*	*

Registration Pass Access

Registration Pass Access - Exhibits

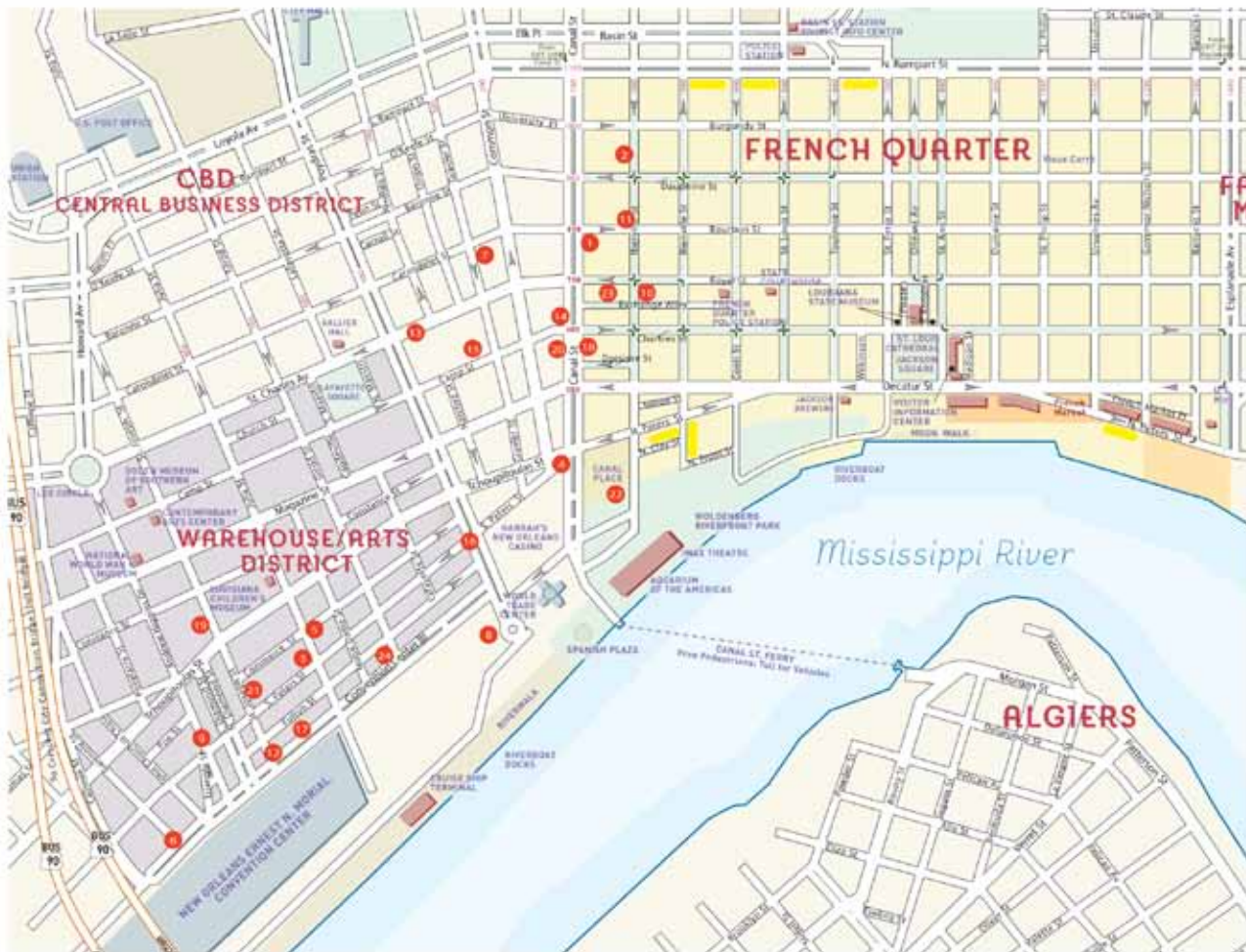
Each registration category provides access to a different set of conference ; summarized below.

Type of Event	Exhibitor 24-hour Access	Exhibit Hall Only Access Tue - Thur
Awards (Thursday)	*	*
Birds-of-a-Feather	*	*
Exhibit Floor	*	*
Exhibitor Forum	*	*
Exhibits Gala Opening (Monday)	*	
Exhibitor's Reception	*	
HPC Matters Plenary (NO BADGE REQUIRED/OPEN TO PUBLIC)	*	*
Invited Talks (Plenary)	*	
Keynote (Tuesday)	*	*
Panels (Friday Only)	*	
Posters	*	*
Student Cluster Competition	*	*

Maps/Daily Schedules

Maps/Daily Schedules

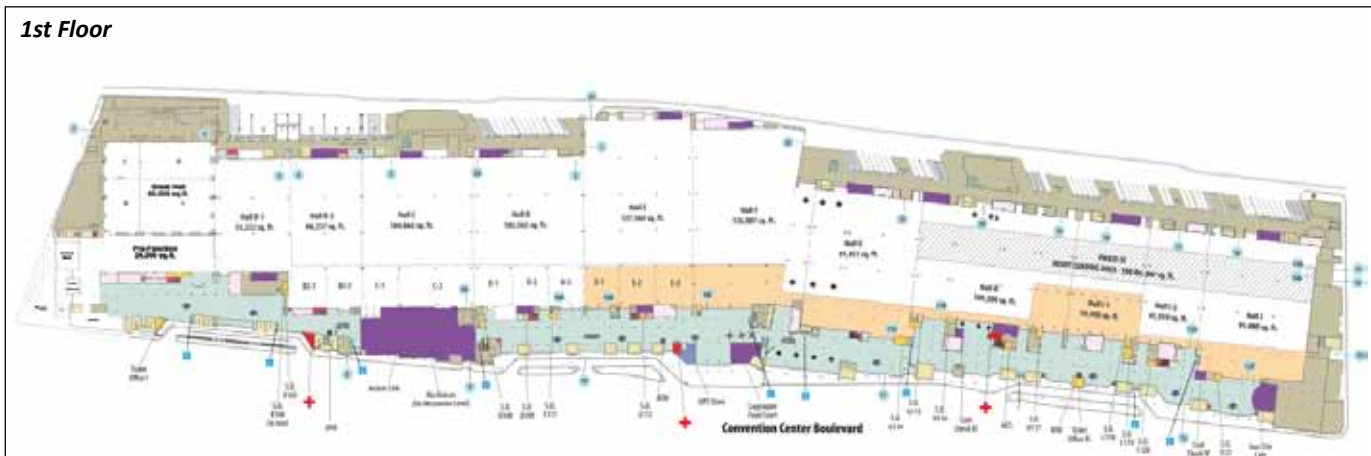
A schedule of each day's activities—by time/event/location—is provided in this section, along with a map of the Downtown area and meeting rooms.



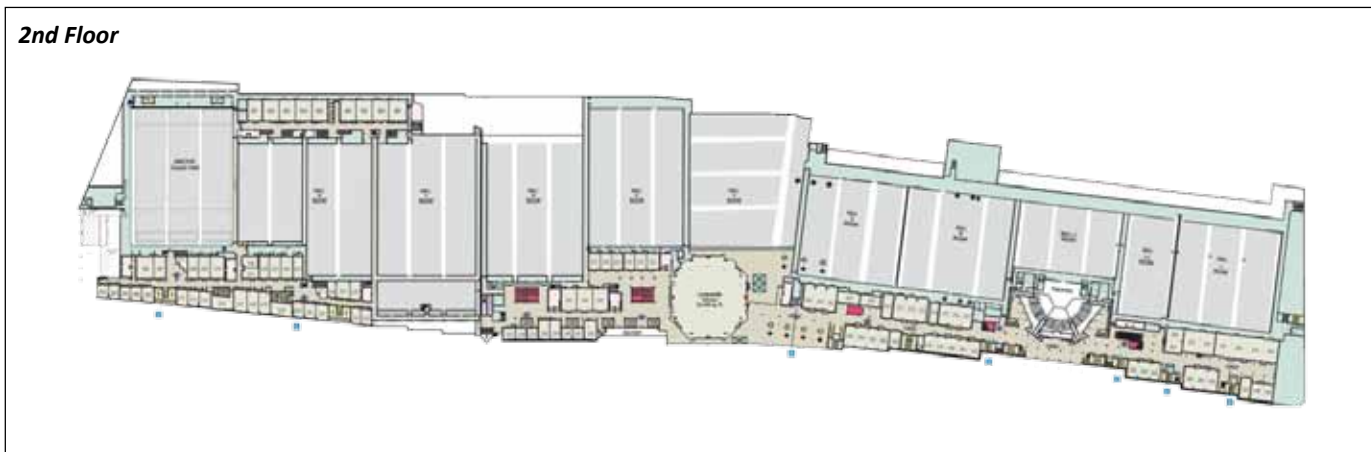
Map #	Hotel	Map #	Hotel
1.	Astor Crowne Plaza	13.	InterContinental New Orleans
2.	Courtyard New Orleans Downtown Iberville	14.	JW Marriott New Orleans
3.	Courtyard by Marriott Convention Center	15.	La Quinta Inn & Suites Downtown
4.	Doubletree Hotel New Orleans	16.	Loews New Orleans Hotel
5.	Embassy Suites Hotel New Orleans Convention Center	17.	New Orleans Downtown Marriott at the Convention Center
6.	Hampton Inn & Suites New Orleans Convention Center	18.	New Orleans Marriott
7.	Hampton Inn Downtown/ French Quarter Area	19.	Residence Inn by Marriott New Orleans - Downtown
8.	Hilton New Orleans Riverside	20.	Sheraton New Orleans Hotel
9.	Hilton Garden Inn New Orleans Convention Center	21.	Springhill Suites by Marriott Convention Center
10.	Hotel Monteleone	22.	Westin New Orleans Canal Place
11.	Hyatt French Quarter	23.	Wyndham French Quarter
12.	Hyatt Place New Orleans Convention Center	24.	Omni Riverfront

Ernest N. Morial Convention Center Floor Plan

1st Floor



2nd Floor



3rd Floor



SATURDAY, NOVEMBER 15

Time	Event Type	Session Title	Location
4pm-5pm	HPCI (BE)	Orientation	288-89
5pm-6:30pm	HPCI (BE)	Submitting Posters: Making Research Known	288-89

SUNDAY, NOVEMBER 16

Time	Event Type	Session Title	Location
8:30am-10am	HPCI (BE)	Plenary I: Introduction to HPC and Its Applications	288-89
8:30am-12pm	Tutorials	A Computation-Driven Introduction to PP in Chapel	388
8:30am-12pm	Tutorials	How to Analyze Performance of Parallel Codes 101	392
8:30am-12pm	Tutorials	MPI-X - Hybrid Programming	393
8:30am-5pm	Tutorials	A "Hands-On" Introduction to OpenMP	395
8:30am-5pm	Tutorials	Parallel Debugging for MPI, Threads, and Beyond	398-99
8:30am-5pm	Tutorials	From "Hello World" to Exascale Using x86, GPUs...	396
8:30am-5pm	Tutorials	Hands-On Practical Hybrid Parallel App Perf Engr	383-84-85
8:30am-5pm	Tutorials	Large Scale Visualization with ParaView	389
8:30am-5pm	Tutorials	Linear Algebra Libraries for HPC...	391
8:30am-5pm	Tutorials	Parallel Computing 101	386-87
8:30am-5pm	Tutorials	Parallel Programming in Modern Fortran	397
8:30am-5pm	Tutorials	Programming the Xeon Phi	394
8:30am-5pm	Tutorials	SciDB -Manage and Analyze TBs of Array Data	390
9am-12:30pm	Workshops	Innovating the Network for Data Intensive Science	274
9am-12:30pm	Workshops	Integrating Comp Science into the Curriculum...	293
9am-5:30pm	Workshops	Big Data Analytics: Challenges and Opportunites	286-87
9am-5:30pm	Workshops	DISCS-2014	298-99
9am-5:30pm	Workshops	E2SC: Energy Efficient Supercomputing	291
9am-5:30pm	Workshops	IA ³ 2014: 4th Workshop on Irregular Applications...	273
9am-5:30pm	Workshops	Many-Task Computing on Clouds, Grids...	295
9am-5:30pm	Workshops	Parallel Data Storage	271-72
9am-5:30pm	Workshops	Sustainable Software for Science: Practice and...	275-76-77
9am-5:30pm	Workshops	Performance Modeling, Benchmarking and Simulation	283-84-85
9am-5:30pm	Workshops	Ultrascale Visualization	294
9am-5:30pm	Workshops	High Performance Computational Finance	292
9am-5:30pm	Workshops	Workflows in Support of Large-Scale Science	297
10am-10:30am	Tutorials	Refreshment Break	Outside 383-399
10am-10:30am	Workshops	Refreshment Break	Outside 283-299
10:30am-11:15am	HPCI (BE)	Session IA: HPC Memory Lane & Future Roadmap (Adv)	288-89
10:30am-11:15am	HPCI (BE)	Session IB: What is HPC, SC and Why Relevant? (Intro)	290

SUNDAY, NOVEMBER 16

Time	Event Type	Session Title	Location
11:15am-12pm	HPCI (BE)	Session IA: Multicore Programming (Adv)	288-89
11:15am-12pm	HPCI (BE)	Session IB: HPC Memory Lane & Future Roadmap (Intro)	290
12pm-1:30pm	Tutorials	Lunch	La Nouvelle
1:30pm-3pm	HPCI (BE)	Session II: HPC Applications and Q&A	288-89
1:30pm-5pm	Tutorials	PGAS & Hybrid MPI+PGAS Programming Models...	393
1:30pm-5pm	Tutorials	Practical Fault Tolerance on Today's SC Systems	388
1:30pm-5pm	Tutorials	Scaling I/O Beyond 100K Cores Using ADIOS	392
1:30pm-5:30pm	Workshops	Education for HPC	293
1:30pm-5:30pm	Workshops	Network-Aware Data Management	274
3pm-3:30pm	Tutorials	Refreshment Break	Outside 383-399
3pm-3:30pm	Workshops	Refreshment Break	Outside 283-299
3:30pm-4:30pm	HPCI (BE)	Session III: Why you Should Love Computational Science	288-89
5:30pm-8pm	HPCI (BE)	Networking Event: Student/Faculty/Professional Posters	288-89-90

MONDAY NOVEMBER 17

Time	Event Type	Session Title	Location
8:30am-10am	HPCI (BE)	Plenary II: Big Data and Exascale Challenges	288-89
8:30am-12pm	Tutorials	In Situ Data Analysis and Vis with ParaView Catalyst	394
8:30am-12pm	Tutorials	InfiniBand and High-Speed Ethernet for Dummies	392
8:30am-12pm	Tutorials	Introducing R: From Your Laptop to HPC and Big Data	391
8:30am-12pm	Tutorials	Parallel Programming with Charm++	393
8:30am-5pm	Tutorials	Advanced MPI Programming	389
8:30am-5pm	Tutorials	Advanced OpenMP: Performance and 4.0 Features	397
8:30am-5pm	Tutorials	Debugging and Perf Tools for MPI and OpenMP 4.0...	398-99
8:30am-5pm	Tutorials	Fault-Tolerance for HPC: Theory and Practice	388
8:30am-5pm	Tutorials	Node-Level Performance Engineering	390
8:30am-5pm	Tutorials	OpenACC: Productive, Portable Perf on Hybrid Systems...	396
8:30am-5pm	Tutorials	OpenCL: A Hands-On Introduction	395
8:30am-5pm	Tutorials	Parallel I/O In Practice	386-87
8:30am-5pm	Tutorials	Python in HPC	383-84-85
9am-5:30pm	Workshops	Python for High Performance & Scientific Computing	291
9am-5:30pm	Workshops	Energy Efficient HPC Working Group	271-72
9am-5:30pm	Workshops	Latest Advances in Scalable Algorithms for Scala	283-84-85
9am-5:30pm	Workshops	ATIP: Japanese Research Toward Next-Gen...	295
9am-5:30pm	Workshops	Co-HPC: Hardware-Software Co-Design for HPC	273
9am-5:30pm	Workshops	ExaMPI14: Exascale MPI 2014	286-87
9am-5:30pm	Workshops	Extreme-Scale Programming Tools	297
9am-5:30pm	Workshops	HP Technical Computers in Dynamic Languages	293

MONDAY, NOVEMBER 17

	Time	Event Type	Session Title	Location
	9am-5:30pm	Workshops	LLVM Compiler Infrastructure in HPC	292
	9am-5:30pm	Workshops	VISTech 2014	294
	9am-5:30pm	Workshops	Accelerator Programming Using Directives	275-76-77
	9am-5:30pm	Workshops	Domain-Specific Languages and High-Level...	298-99
	10am-10:30am	Tutorials	Refreshment Break	Outside 383-399
	10am-10:30am	Workshops	Refreshment Break	Outside 283-299
	10:30am-12pm	HPCI (BE)	Session IVA: Programming for Exascale (Advanced)	288-89
	10:30am-12pm	HPCI (BE)	Session IVB: Parallel Programming (Introduction)	290
	12pm-1:30pm	Tutorials	Lunch	La Nouvelle
	1:30pm-3:30pm	HPCI (BE)	Mentor/Protégé Session and Mixer	288-89
	1:30pm-5pm	Tutorials	Designing/Using High-End Systems with InfiniBand...	392
	1:30pm-5pm	Tutorials	Effective HPC Visualization Using VisIt	391
	1:30pm-5pm	Tutorials	Enhanced Campus Bridging Using Globus and...	394
	1:30pm-5pm	Tutorials	Introductory and Advanced OpenSHMEM Programming	393
	3pm-3:30pm	Tutorials	Refreshment Break	Outside 383-399
	3pm-3:30pm	Workshops	Refreshment Break	Outside 283-299
	3pm-5pm	HPCI (Undergraduates)	Experiencing HPC for Undergraduates Orientation	290
	3:30pm-7pm	HPCI (BE)	Programming Challenge	288-89
	5pm-5:30pm	HPCI	First Time at SC? Here's How to Get Oriented!	274
	5:30pm-6:45pm	Plenary Talk	SC14 HPC Matters Plenary	N.O. Theater
	7pm-9pm	HPC IS/ET	Emerging Technologies Exhibits	Show Floor-233
	7:15pm-8:30pm	HPC IS/ET	Gala HPC Impact Showcase	Show Floor-233
	7pm-9pm	Social Event	Gala Opening Reception	Exhibit Halls G-J
	7:30pm-9pm	HPCI (SCC)	Student Cluster Competition Kickoff	Show Floor-410

TUESDAY, NOVEMBER 18

	Time	Event Type	Session Title	Location
	8:30am-10am	Keynote	Brian Greene, Professor of Physics & Mathematics, Columbia University	N.O. Theater
	8:30am-5pm	Posters	Poster Display (Including ACM SRC)	N.O. Theater Lobby
	10am-10:30am	Technical Program	Refreshment Break	Outside of 383-399
	10am-6pm	HPC IS/ET	Emerging Technologies Exhibits	Show Floor-233
	10am-6pm	HPCI (SCC)	Student Cluster Competition	Show Floor-410
	10:20am-2:20pm	HPC IS/ET	HPC Impact Showcase Presentations	Show Floor-233
	10:30am-12pm	Exhibitor Forum	Hardware and Architecture	292
	10:30am-12pm	Exhibitor Forum	Moving, Managing and Storing Data	291
	10:30am-12pm	HPCI (Undergraduates)	Introduction to HPC Research	295
	10:30am-12pm	Invited Speakers	Invited Talks (Masoud Mohseni; David Abramson)	N.O. Theater
	10:30am-12pm	Panels	Funding Strategies for HPC Software Beyond Borders	383-84-85
	10:30am-12pm	Papers	Heterogeneity and Scaling in Applications	393-94-95
	10:30am-12pm	Papers	Memory and Microarchitecture	391-92
	10:30am-12pm	Papers	Performance Measurement	388-89-90
	12:15pm-1:15pm	Birds of a Feather	Asynchronous Many-Task Programming Models...	396
	12:15pm-1:15pm	Birds of a Feather	Chapel Lightning Talks 2014	293
	12:15pm-1:15pm	Birds of a Feather	Code Optimization War Stories	298-99
	12:15pm-1:15pm	Birds of a Feather	Experimental Infrastructures for Open Cloud Research	388-89-90
	12:15pm-1:15pm	Birds of a Feather	HDF5: State of the Union	398-99
	12:15pm-1:15pm	Birds of a Feather	HPC Sys and Data Center Energy Efficiency Metrics...	294
	12:15pm-1:15pm	Birds of a Feather	LLVM in HPC: Uses and Desires	283-84-85
	12:15pm-1:15pm	Birds of a Feather	Lustre Community: At the Heart of HPC and Big Data	275-76-77
	12:15pm-1:15pm	Birds of a Feather	Ninth Graph500 List	286-87
	12:15pm-1:15pm	Birds of a Feather	Operating System and Run-Time for Exascale	391-92
	12:15pm-1:15pm	Birds of a Feather	Quantum Chemistry 500	295
	12:15pm-1:15pm	Birds of a Feather	Resilience & Power-Efficiency Challenges at Exascale..	292
	12:15pm-1:15pm	Birds of a Feather	SAGE for Global Collaboration	393-94-95
	12:15pm-1:15pm	Birds of a Feather	SIGHPC Annual Meeting	383-84-85
	12:15pm-1:15pm	Birds of a Feather	The 2014 HPC Challenge Awards	273
	12:15pm-1:15pm	Birds of a Feather	Towards European Leadership in HPC	386-87
	1:30pm-2:15pm	Awards	SC14 Test of Time Award-Special Lectures	386-87
	1:30pm-3pm	Awards	ACM Gordon Bell Finalists	N.O. Theater
	1:30pm-3pm	Exhibitor Forum	Moving, Managing and Storing Data	291
	1:30pm-3pm	Exhibitor Forum	Hardware and Architecture	292
	1:30pm-3pm	HPCI (BE)	Interviewing & Mock Interviews	288-89
	1:30pm-3pm	Panels	Changing Op Sys is Harder than Changing Prog Lang	383-84-85
	1:30pm-3pm	Papers	Accelerators	388-89-90
	1:30pm-3pm	Papers	Best Practices in File Systems	393-94-95
	1:30pm-3pm	Papers	Earth and Space Sciences	391-92
	3pm-3:30pm	Technical Program	Refreshment Break	Outside 383-399

TUESDAY, NOVEMBER 18

	Time	Event Type	Session Title	Location
	3:30pm-5pm	Exhibitor Forum	Moving, Managing and Storing Data	291
	3:30pm-5pm	Exhibitor Forum	Hardware and Architecture	292
	3:30pm-5pm	Invited Speakers	Invited Talks (Horst Simon; Valerie Taylor)	N.O. Theater
	3:30pm-5pm	Panels	HPC Productivity or Performance: Choose One	383-84-85
	3:30pm-5pm	Papers	Compiler Analysis and Optimization	393-94-95
	3:30pm-5pm	Papers	Networks	388-89-90
	3:30pm-5pm	Papers	Parallel Algorithms	391-92
	4:30pm-6pm	HPC IS/ET	Emerging Technologies Exhibits	Show Floor-233
	5pm-7pm	HPCI (BE)	Guided Tours of the Poster Session	N.O.Theater Lobby
	5:15pm-7pm	Posters	Reception, Posters (Including ACM SRC)	N.O.Theater Lobby
	5:30pm-7pm	Birds of a Feather	Compiler Vec & Achieving Effective SIMD...	391-92
	5:30pm-7pm	Birds of a Feather	Design, Commissioning, Controls for Liquid Cooling...	383-84-85
	5:30pm-7pm	Birds of a Feather	Forming a Support Org for HPC User Serv Providers	297
	5:30pm-7pm	Birds of a Feather	Publishing Computational Data	286-87
	5:30pm-7pm	Birds of a Feather	Here Come OpenMP 4 Implementations & OpenMP	291
	5:30pm-7pm	Birds of a Feather	How TORQUE is Changing to Meet New Demands	293
	5:30pm-7pm	Birds of a Feather	HPC Systems Engr, Suffering, & Administration	294
	5:30pm-7pm	Birds of a Feather	Integrated Optimization of Perf, Power & Resilience...	298-99
	5:30pm-7pm	Birds of a Feather	Joint NSF/ENG and AFOSR Initiative on DDS	283-84-85
	5:30pm-7pm	Birds of a Feather	Migration to Lustre 2.5 and Beyond...	292
	5:30pm-7pm	Birds of a Feather	MPICH: An HP Open-Source MPI Implementation	386-87
	5:30pm-7pm	Birds of a Feather	OpenCL: Version 2.0 and Beyond	275-76-77
	5:30pm-7pm	Birds of a Feather	OpenPower: Future Directions with HPC	273
	5:30pm-7pm	Birds of a Feather	Perf Analysis and Sim of MPI Apps & Runtimes...	398-99
	5:30pm-7pm	Birds of a Feather	PRObE - 1500 New Machines for Systems Research	396
	5:30pm-7pm	Birds of a Feather	Reconfigurable Supercomputing	393-94-95
	5:30pm-7pm	Birds of a Feather	Strategies for Academic HPC Centers	388-89-90
	5:30pm-7pm	Birds of a Feather	TOP500 Supercomputers	N. O. Theater
	5:30pm-7pm	Birds of a Feather	Women in HPC: Mentorship and Leadership	295

WEDNESDAY, NOVEMBER 19

	Time	Event Type	Session Title	Location
	8:30am-5pm	Posters	Poster Display (Including ACM SRC)	N.O.Theater Lobby
	8:30am-10am	Awards	Cray/Fernbach/Kennedy Award Recipient Talks	N.O. Theater
	10am-10:30am	Technical Program	Refreshment Break	Outside 383-399
	10am-3pm	HPCI	Student-Postdoc Job & Opportunity Fair	288-89
	10am-5pm	HPCI (SCC)	Student Cluster Competition	Show Floor-410
	10am-6pm	HPC IS/ET	Emerging Technologies Exhibits	Show Floor-233
	10:20am-2:20pm	HPC IS/ET	HPC Impact Showcase Presentations	Show Floor-233
	10:30am-12pm	Exhibitor Forum	Software for HPC	291
	10:30am-12pm	Exhibitor Forum	HPC Futures and Exascale	292
	10:30am-12pm	HPCI (Undergraduates)	HPC For Undergraduates: Grad Student Perspective	295
	10:30am-12pm	Invited Speakers	Invited Talks (Michael Heath; Jim Sexton)	N.O. Theater
	10:30am-12pm	Panels	Analyst Crossfire & Advanced Cyberinfrastructure	383-84-85
	10:30am-12pm	Papers	MPI	393-94-95
	10:30am-12pm	Papers	Big Data Analysis	391-92
	10:30am-12pm	Papers	High Performance Genomics	388-89-90
	10:30am-12pm	Vis & Data Analytics	Showcase	386-87
	12:15pm-1:15pm	Birds of a Feather	App Readiness & Portability for Leadership Computing	293
	12:15pm-1:15pm	Birds of a Feather	Dynamic Power Mgmt for MW-Sized SC Centers	292
	12:15pm-1:15pm	Birds of a Feather	Getting Scientific Software Installed: Tools and...	298-99
	12:15pm-1:15pm	Birds of a Feather	HPGMG Proposal for New Top500 Benchmark	294
	12:15pm-1:15pm	Birds of a Feather	HPC Debugging Techniques	393-94-95
	12:15pm-1:15pm	Birds of a Feather	New SIGHPC Education Chapter Meeting	295
	12:15pm-1:15pm	Birds of a Feather	HPC Job Scheduling Challenges & Successes from...	291
	12:15pm-1:15pm	Birds of a Feather	Open MPI State of the Union	283-84-85
	12:15pm-1:15pm	Birds of a Feather	PGAS	273
	12:15pm-1:15pm	Birds of a Feather	Power API for HPC: Standardizing Power Measurements...	286-87
	12:15pm-1:15pm	Birds of a Feather	Programming Abstractions for Data Locality	391-92
	12:15pm-1:15pm	Birds of a Feather	Python for HP and Scientific Computing	275-76-77
	12:15pm-1:15pm	Birds of a Feather	SC Community Town Hall	396
	12:15pm-1:15pm	Birds of a Feather	Scalable Industrial Software: Future or Fantasy?	386-87
	12:15pm-1:15pm	Birds of a Feather	Open Science Data Cloud & PIRE Fellowships...	297
	12:15pm-1:15pm	Birds of a Feather	The OpenPOWER Foundation: Addressing HPC...	383-84-85
	1:30pm-3pm	Awards	ACM Gordon Bell Finalists	N.O. Theater
	1:30pm-3pm	Exhibitor Forum	Storage for HPC	291
	1:30pm-3pm	Exhibitor Forum	HPC Futures and Exascale	292
	1:30pm-3pm	Panels	Can We Avoid Bldg an Exascale "Stunt" Machine?	383-84-85
	1:30pm-3pm	Papers	Cloud Computing I	391-92
	1:30pm-3pm	Papers	Graph Algorithms	388-89-90
	1:30pm-3pm	Papers	Hardware Vulnerability and Recovery	393-94-95

WEDNESDAY, NOVEMBER 19

	Time	Event Type	Session Title	Location
	2:30pm-6pm	HPC IS/ET	Emerging Technologies Presentations	Show Floor-233
	3pm-3:30pm	Technical Program	Refreshment Break	Outside 383-99
	3:30pm-5pm	Exhibitor Forum	Software for HPC	291
	3:30pm-5pm	Exhibitor Forum	HPC Futures and Exascale	292
	3:30pm-5pm	Invited Speakers	Invited Talks (Larry Smarr; Cleve Moler)	N.O. Theater
	3:30pm-5pm	Panels	Future of Mem Tech for Exascale & Beyond II	383-84-85
	3:30pm-5pm	Papers	I/O and Dynamic Optimization	391-92
	3:30pm-5pm	Papers	Quantum Simulations in Materials and Chemistry	388-89-90
	3:30pm-5pm	Papers	Resilience	393-94-95
	3:30pm-5pm	ACM SRC Posters	ACM SRC Posters Presentations	386-87
	5pm-6pm	HPCI (SCC)	Student Cluster Competition Grand Finale	Show Floor-410
	5:30pm-7pm	Birds of a Feather	Application Experiences with Emerging PGAS APIs	386-87
	5:30pm-7pm	Birds of a Feather	Challenges in Managing Small HPC Centers	295
	5:30pm-7pm	Birds of a Feather	Chapel Users Group Meeting	383-84-85
	5:30pm-7pm	Birds of a Feather	Codesign for DOE's Comp Science Community	292
	5:30pm-7pm	Birds of a Feather	Developing HPC Workforce of the Future	393-94-95
	5:30pm-7pm	Birds of a Feather	Evaluation Infrastructure for Future Large-Scale Chips...	391-92
	5:30pm-7pm	Birds of a Feather	Experiences Delivering Campus Data Service w/Globus	298-99
	5:30pm-7pm	Birds of a Feather	Integrating HPC Technologies in BigData Architectures	283-84-85
	5:30pm-7pm	Birds of a Feather	Monitoring Large-Scale HPC Systems...	297
	5:30pm-7pm	Birds of a Feather	National Academies Study on Future Directions for NSF ACI to Support U.S. Science in 2017-2020	396
	5:30pm-7pm	Birds of a Feather	OpenACC API User Experience, Vendor Reaction...	275-76-77
	5:30pm-7pm	Birds of a Feather	OpenSHMEM: Further Develop the SHMEM Standard...	294
	5:30pm-7pm	Birds of a Feather	Programmable Storage Systems	398-99
	5:30pm-7pm	Birds of a Feather	Super-R: Supercomputing and R for Data-Intensive...	388-89-90
	5:30pm-7pm	Birds of a Feather	The Eclipse Parallel Tools Platform	291
	5:30pm-7pm	Birds of a Feather	The HPCG Benchmark...	273
	5:30pm-7pm	Birds of a Feather	MPI 3.1 and Plans for MPI 4.0	293
	5:30pm-7pm	Birds of a Feather	Xeon Phi Users Group	286-87

THURSDAY, NOVEMBER 20

	Time	Event Type	Session Title	Location
	8:30am-10am	Invited Speaker	Invited Plenary Talks (Philip Bourne; Lincoln Wallen)	N. O.Theater
	8:30am-5pm	Posters	Poster Display (Including ACM SRC)	N.O.Theater Lobby
	10am-10:30am	Technical Program	Refreshment Break	Outside 383-399
	10am-3pm	HPC IS/ET	Emerging Technologies Exhibits	Show Floor-233
	10:20am-1pm	HPC IS/ET	HPC Impact Showcase Presentations	Show Floor-233
	10:30am-12pm	Doctoral Showcase	Dissertation Research	386-87
	10:30am-12pm	Exhibitor Forum	Effective Application of HPC	291
	10:30am-12pm	Exhibitor Forum	Moving, Managing and Storing Data	292
	10:30am-12pm	HPCI (Undergraduates)	Careers in HPC	295
	10:30am-12pm	Invited Speakers	Invited Talks (Dona Crawford; Mateo Valero)	N.O. Theater
	10:30am-12pm	Panels	Multi-Sys Supercomputer Procurements – Good Idea?	383-84-85
	10:30am-12pm	Papers	Machine Learning and Data Analytics	388-89-90
	10:30am-12pm	Papers	Numerical Kernels	391-92
	10:30am-12pm	Papers	Power and Energy Efficiency	393-94-95
	12pm-12:45pm	HPCI	HPCI Scavenger Hunt Awards	274
	12:15pm-1:15pm	Birds of a Feather	Analyzing Parallel I/O	286-87
	12:15pm-1:15pm	Birds of a Feather	Charm++: Adaptive & Asynchronous Parallel Prog	291
	12:15pm-1:15pm	Birds of a Feather	Comp & Data Challenges in Genomic Sequencing	273
	12:15pm-1:15pm	Birds of a Feather	Config Mgmt Tools for Large Scale HPC Clusters	388-89-90
	12:15pm-1:15pm	Birds of a Feather	Defining Interfaces for Interoperable Simulations...	295
	12:15pm-1:15pm	Birds of a Feather	Understanding User-Level Activity on Today's Supercomputers w/XALT	386-87
	12:15pm-1:15pm	Birds of a Feather	Evolving HPC Software Ecosystem for Emerging...	292
	12:15pm-1:15pm	Birds of a Feather	HPC Centers for Global Health	398-99
	12:15pm-1:15pm	Birds of a Feather	Managing HPC Centers: Communicating the Value of...	393-94-95
	12:15pm-1:15pm	Birds of a Feather	OpenStack Tech & HPC Cloud Solutions	297
	12:15pm-1:15pm	Birds of a Feather	Propose & Define an Abstracted Interface for Gen Usage	396
	12:15pm-1:15pm	Birds of a Feather	Slurm User Group	283-84-85
	12:15pm-1:15pm	Birds of a Feather	System Testing & Resiliency in HPC	391-92
	12:15pm-1:15pm	Birds of a Feather	The Future of Fortran	293
	12:15pm-1:15pm	Birds of a Feather	The Green500 List and its Continuing Evolution	275-76-77
	12:15pm-1:15pm	Birds of a Feather	OCR Framework for Extreme Scale Systems	294
	12:15pm-1:15pm	Birds of a Feather	The Square Kilometer Array: Next Gen Sensor Networks...	383-84-85
	12:30pm-1:30pm	Awards	Award Presentations	N.O. Theater
	1:30pm-3pm	Doctoral Showcase	Dissertation Research	386-87
	1:30pm-3pm	Exhibitor Forum	Effective Application of HPC	291
	1:30pm-3pm	Exhibitor Forum	Hardware and Architecture	292
	1:30pm-3pm	HPCI (SCC)	Get Involved in the Student Cluster Competition	295
	1:30pm-3pm	Panels	InfoSymbioticSystems/DDDAS	383-84-85
	1:30pm-3pm	Papers	Data Locality and Load Balancing	388-89-90
	1:30pm-3pm	Papers	Optimized Checkpointing	393-94-95

THURSDAY, NOVEMBER 20

Time	Event Type	Session Title	Location
1:30pm-3pm	Papers	Sparse Solvers	391-92
2pm-3pm	HPC IS/ET	Emerging Technologies Presentations	Show Floor-233
3pm-3:30pm	Technical Program	Refreshment Break	Outside 383-399
3:30pm-5pm	Doctoral Showcase	Dissertation Research	386-87
3:30pm-5pm	Exhibitor Forum	Effective Applications of HPC	291
3:30pm-5pm	Exhibitor Forum	Hardware and Architecture	292
3:30pm-5pm	HPCI (BE)	Wrap Up	295
3:30pm-5pm	Panels	Challenges with Liquid Cooling...	383-84-85
3:30pm-5pm	Papers	Cloud Computing II	388-89-90
3:30pm-5pm	Papers	Large-Scale Visualization	391-92
3:30pm-5pm	Papers	Memory System Energy Efficiency	393-94-95
6pm-9pm	Social Event	Technical Program Conference Reception (Shuttle transportation	The Sugar Mill

between the hotels and convention center will run until 9pm from Lobby I.)



FRIDAY, NOVEMBER 21

	Time	Event Type	Session Title	Location
	8:30am-10am	Panels	Beyond Von Neumann, Neuromorphic Sys & Architecture	391-92
	8:30am-10am	Panels	Return of HPC Survivor: Outwit, Outlast, Outcompute	383-84-85
	8:30am-12pm	Workshops	Data Intensive Computing in the Clouds	292
	8:30am-12pm	Workshops	Gateway Computing Environments	273
	8:30am-12pm	Workshops	HPC User Support Tools	297
	8:30am-12pm	Workshops	Software Engineering for HPC in CSE	283-84-85
	8:30am-12pm	Workshops	Visual Performance Analysis	298-99
	8:30am-12pm	Workshops	Women in HPC	275-76-77
	8:30am-12pm	Workshops	Best Practices for HPC Training	294
	10am-10:30am	Technical Program	Refreshment Break	Outside 383-399
	10am-10:30am	Workshops	Refreshment Break	Outside 283-299
	10:30am-12pm	Panels	ROI from Academic Supercomputing	391-92
	10:30am-12pm	Panels	The Roadblock Ahead: Power Usage for Storage & I/O	383-84-85



Award Talks/ Award Presentations

Award Talks/Award Presentations

Each year, SC showcases not only the best and brightest stars of high-performance computing, but also its rising stars and those who have made a lasting impression. SC Awards provide one way these individuals are recognized. This page describes the various awards given each year at the SC Conference.

IEEE-CS Seymour Cray Computer Engineering Award

The Seymour Cray Computer Engineering Award recognizes innovative contributions to HPC systems that best exemplify the creative spirit of Seymour Cray. Nominations that recognize the design and engineering and intellectual leadership in creating innovative and successful HPC systems are especially solicited. The award includes a \$10,000 honorarium.

IEEE-CS Sidney Fernbach Memorial Award

The Sidney Fernbach Memorial Award honors innovative uses of HPC in problem solving. Nominations that recognize creation of widely used and innovative software packages, application software, and tools are especially solicited. A certificate and a \$2000 honorarium are given to the winner.

ACM/IEEE-CS Ken Kennedy Award

The Ken Kennedy Award recognizes substantial contributions to programmability and productivity in computing and substantial community service or mentoring contributions. The award honors the remarkable research, service, and mentoring contributions of Ken Kennedy and includes a \$5000 honorarium.

ACM Gordon Bell Prize

The Gordon Bell Prize is awarded each year to recognize outstanding achievement in high-performance computing. The award tracks the progress over time of parallel computing, with particular emphasis on rewarding innovation in applying high-performance computing to applications in science, engineering, and large-scale data analytics. Prizes may be awarded for peak performance or special achievements in scalability and time-to-solution on important science and engineering problems. The award is administered by ACM, and includes a \$10,000 cash prize. ACM/IEEE-CS George Michael Memorial HPC Ph.D. Fellowship In memory of George Michael, one of the founding fathers of

the SC Conference series, this award honors exceptional PhD students throughout the world whose research focus is on high-performance computing applications, networking, storage, or large-scale data analysis using the most powerful computers that are currently available. Fellowship winners will be selected based on overall potential for research excellence, the degree to which technical interests align with those of the HPC community, academic progress to date, recommendations by their advisors, and a demonstration of current and anticipated use of HPC resources

Best Paper

Each year, the SC Technical Papers Committee identifies one paper as the best paper from the Conference's Technical Program. The award will be announced at the SC14 Awards Ceremony.

Best Student Paper

Each year, the SC Technical Papers Committee identifies one paper as the best paper written primarily by a student or students and presented in the Conference's Technical Program. Papers must be identified as student papers when submitted to be eligible for this award. This award will be announced at the SC14 Awards Ceremony.

Best Poster

Each year, the SC Posters Committee identifies one poster as the best poster presented as part of the Conference's Technical Program. The award will be announced at the SC14 Awards Ceremony.

Test of Time Award

The "Test of Time" award recognizes a paper from a past conference that has deeply influenced the HPC discipline. It is a mark of historical impact, and requires clear evidence that the paper has changed HPC trends. The award was originally presented at SC13, as part of the celebration of the SC 25th anniversary. It is presented annually to a single paper, selected from the conference proceedings of 10-25 years ago.

ACM Student Research Competition

SC hosts an ACM Student Research Competition as part of its research poster activities within the Technical Program. Up to three awards are given in both undergraduate and graduate student categories. The award will be announced at the SC14 Awards Ceremony.

Award Talks/Award Presentations

Tuesday, November 18

ACM Gordon Bell Finalists

Chair: Taisuke Boku (University of Tsukuba)

1:30pm-3pm

Room: New Orleans Theater

Petascale High Order Dynamic Rupture Earthquake Simulations on Heterogeneous Supercomputers

Alexander Heinecke (Technical University of Munich/Intel Corporation), Alexander Breuer, Sebastian Rettenberger, Michael Bader (Technical University of Munich), Alice-Agnes Gabriel, Christian Pelties (Ludwig Maximilian University of Munich), Arndt Bode (Leibniz Supercomputing Center/Technical University of Munich), William Barth (Texas Advanced Computing Center), Xiang-Ke Liao (National University of Defense Technology), Karthikeyan Vaidyanathan, Mikhail Smelyanskiy, Pradeep Dubey (Intel Corporation)

We present an end-to-end optimization of the innovative Arbitrary high-order DERivative Discontinuous Galerkin (ADER-DG) software SeisSol targeting Intel® Xeon Phi™ coprocessor platforms, achieving unprecedented earthquake model complexity through coupled simulation of full frictional sliding and seismic wave propagation. SeisSol exploits unstructured meshes to flexibly adapt for complicated geometries in realistic geological models. Seismic wave propagation is solved simultaneously with earthquake faulting in a multi-physical manner leading to a heterogeneous solver structure. Our architecture-aware optimizations deliver up to 50% of peak performance, and introduce an efficient compute-communication overlapping scheme shadowing the multiphysics computations. SeisSol delivers near-optimal weak scaling, reaching 8.6 DP-PFLOPS on 8,192 nodes of the Tianhe-2 supercomputer. Our performance model projects reaching 18-20 DP-PFLOPS on the full Tianhe-2 machine. Of special relevance to modern civil engineering needs, our pioneering simulation of the 1992 Landers earthquake shows highly detailed rupture evolution and ground motion at frequencies up to 10 Hz.

Physics-Based Urban Earthquake Simulation Enhanced by 10.7 BlnDOF x 30 K Time-Step Unstructured FE Non-Linear Seismic Wave Simulation

Tsuyoshi Ichimura (Earthquake Research Institute and Department of Civil Engineering, University of Tokyo), Kohei Fujita (RIKEN Advanced Institute for Computational Science & Japan Society for the Promotion of Science), Seizo Tanaka (Earthquake Research Institute and Department of Civil Engineering, University of Tokyo), Muneo Hori (Earthquake Research Institute and Department of Civil Engineering, University of Tokyo & RIKEN Advanced Institute for Computational Science), Maddegadara Lalith (Earthquake Research Institute and Department of Civil Engineering, University of Tokyo), Yoshihisa Shizawa (Research Organization for Information Science and Technology), Hiroshi Kobayashi (Research Organization for Information Science and Technology)

With the aim of dramatically improving the reliability of urban earthquake response analyses, we developed an unstructured 3-D finite-element-based MPI-OpenMP hybrid seismic wave amplification simulation code, GAMERA. On the K computer, GAMERA was able to achieve a size-up efficiency of 87.1% up to the full K computer. Next, we applied GAMERA to a physics-based urban earthquake response analysis for Tokyo. Using 294,912 CPU cores of the K computer for 11 h, 32 min, we analyzed the 3-D non-linear ground motion of a 10.7 BlnDOF problem with 30 K time steps. Finally, we analyzed the stochastic response of 13,275 building structures in the domain considering uncertainty in structural parameters using 3 h, 56 min of 80,000 CPU cores of the K computer. Although a large amount of computer resources is needed presently, such analyses can change the quality of disaster estimations and are expected to become standard in the future.

Real-Time Scalable Cortical Computing at 46 Giga-Synaptic OPS/Watt with ~100x Speedup in Time-to-Solution and ~100,000x Reduction in Energy-to-Solution

Andrew S. Cassidy, Rodrigo Alvarez-Icaza, Filipp Akopyan, Jun Sawada, John V. Arthur, Paul A. Merolla, Pallab Datta, Marc Gonzalez Tallada, Brian Taba, Alexander Andreopoulos, Arnon Amir, Steven K. Esser, Jeff Kusnitz, Rathinakumar Appuswamy, Chuck Haymes, Bernard Brezzo, Roger Moussalli, Ralph Bellofatto, Christian Baks, Michael Mastro, Kai Schleupen, Charles E. Cox, Ken Inoue, Steve Millman, Nabil Imam, Emmett McQuinn, Yutaka T. Nakamura, Ivan Vo, Chen Guo, Don Nguyen, Scott Lekuch, Sameh Assa, Daniel Friedman, Bryan L. Jackson, Myron D. Flickner, William P. Risk (IBM), Rajit Manohar (Cornell University), Dharmendra S. Modha (IBM)

Drawing on neuroscience, we have developed a parallel, event-driven kernel for neurosynaptic computation that is efficient with respect to computation, memory, and communication. Building on the previously demonstrated highly-optimized software expression of the kernel, here, we demonstrate TrueNorth, a co-designed silicon expression of the kernel. TrueNorth achieves five orders of magnitude reduction in energy-to-solution and two orders of magnitude speedup in time-to-solution, when running computer vision applications and complex recurrent neural network simulations. Breaking path with the von Neumann architecture, TrueNorth is a 4,096 core, 1 million neuron, and 256 million synapse brain-inspired neurosynaptic processor, that consumes 65mW of power running at real-time and delivers performance of 46 Giga-Synaptic OPS/Watt. We demonstrate seamless tiling of TrueNorth chips into arrays, forming a foundation for cortex-like scalability. TrueNorth's unprecedented time-to-solution, energy-to-solution, size, scalability, and performance combined with the underlying flexibility of the kernel enable a broad range of cognitive applications.

SC14 Test of Time Award Special Lectures

Chair: Ewing Lusk (Argonne National Laboratory)

1:30pm-2:15pm

Room: 386-87

Graphs and HPC in the '90s: A Distant Mirror or Merely Distant?

Bruce Hendrickson, Rob Leland (Sandia National Laboratories)

Bio: Bruce Hendrickson is Senior Manager for Extreme Scale Computing at Sandia National Laboratories and Affiliated Professor of Computer Science at the University of New Mexico. He is also a Fellow of the Society for Industrial and Applied Mathematics. Bruce has worked on math, algorithms and software for a wide range of parallel computing and scientific applications. Bruce received his PhD in Computer Science from

Cornell, after obtaining degrees in Mathematics and Physics from Brown. He is the author of around 100 scientific papers, has served on the editorial boards of a range of journals in parallel and scientific computing, and has helped organize a number of international meetings. His research interests include combinatorial scientific computing, parallel algorithms, linear algebra, graph algorithms, scientific software, data mining and computer architecture.

Bio: Rob Leland studied undergraduate electrical engineering at Michigan State University. He attended Oxford University where he studied applied mathematics and computer science, completing a Ph.D. in Parallel Computing in 1989. He joined the Parallel Computing Sciences Department at Sandia National Laboratories in 1990 and pursued work principally in parallel algorithm development, sparse iterative methods and applied graph theory. There he co-authored Chaco, a graph partitioning and sequencing toolkit widely used to optimize parallel computations. In 1995 he served as a White House Fellow advising the Deputy Secretary of the Treasury on technology modernization at the IRS. Upon returning to Sandia he led the Parallel Computing Sciences Department and the Computer and Software Systems Research Group. In 2005 he became Director of the Computing and Networking Services at Sandia, with responsibility for production computing platforms, voice and data networks, desktop support and cyber security for the laboratory. Since March of 2010 he has served as Director of Computing Research, leading a vertically integrated set of capabilities spanning computer architecture, math and computing science, algorithm and tool development, computational sciences and cognitive sciences. Rob also served during this period as Director of Sandia's Climate Security Program which is focused on helping the nation understand and prepare for the national security impacts of climate change. Most recently, Rob has been working at the White House in the Office of Science and Technology Policy leading development of a new national strategy for High Performance Computing.

Abstract: The early 1990s was a time of tremendous change and uncertainty in high performance computing. Multiple visions for future architectures and programming models were competing for market-share, and key abstractions and algorithms were still being formulated. Today, after nearly 20 years of remarkable stability in core HPC concepts, we are entering another era of uncertainty. This talk will explore lessons from that long-ago era that might be relevant for the challenges of today. More specifically, the talk will recount the role of graph theoretic abstractions in helping to mask the complexity of underlying hardware (as highlighted in an impactful paper from SC'95). Looking forward, how can we protect application scientists from the vastly greater complexity of upcoming platforms?

Wednesday, November 19

Cray/Fernbach/Kennedy Award Recipients Talks

Chair: Barbara Chapman (University of Houston)

8:30am-10am

Room: New Orleans Theater

A Half-century, Brief, Personalized History of Supercomputing

Gordon Bell (Microsoft)

Bio: Gordon Bell is a Researcher Emeritus at Microsoft's Silicon Valley Laboratory. He has worked as a computer designer, parallel architecture researcher, administrator, supporter, funder, critic, and historian. He spent 23 years at Digital Equipment Corporation (part of HP) as Vice President of Research and Development. He was the architect of various mini- and time-sharing computers and led the development of DEC's VAX. Bell has been involved in, or responsible for, the design of products at Digital, Encore, Ardent, and a score of other companies. Bell has a BS and MS degree from MIT (1956-57), D. Eng. (hon.) from WPI (1993) and D.Sci and Tech (hon.) from CMU (2010). During 1966-72 he was Professor of Computer Science and Electrical Engineering at Carnegie Mellon University. In 1986-1987 he was the first Assistant Director of the National Science Foundation's Computing Directorate, CISE including the Supercomputing Centers program. He led the National Research and Education (NREN) Network panel that became the Internet and was an author of the first High Performance Computer and Communications Initiative.

Bell has authored books and papers about computer structures, lifelogging, and startup companies. In April 1991, Addison-Wesley published *High Tech Ventures: The Guide to Entrepreneurial Success*. In 2009, Dutton published *Total Recall* that describes the journey to storing one's entire life that was written with Jim Gemmell. He is a founder and board member of the Computer History Museum, Mountain View, California.

Bell is a member of various professional organizations including the American Academy of Arts and Sciences (Fellow), American Association for the Advancement of Science (Fellow), ACM (Fellow), IEEE (Fellow and Computer Pioneer), and the National Academy of Engineering, National Academy of Science and the Australian Academy of Technological Science and Engineering. His awards include: IEEE's McDowell and Eckert-Mauchly Awards, the Von Neumann Medal, and the 1991 National Medal of Technology.

Mr. Bell lives in San Francisco and Sydney, Australia with his wife, Sheridan Sinclair-Bell.

Abstract: On visiting Livermore and seeing LARC a half century ago, I marveled at a kind of large computer built at the edge of technology and a market's ability to pay. As a machine builder, researcher, administrator, supporter, funder, critic, and historian there's nothing quite like supercomputers. I'll posit the mono-memory aka SR Cray computer era and multi-computer era consistent with the 3 era Livermore overview e.g. Rob Neely. The decade (mid-80s) of the search for performance and scaling where about 40 companies lost their lives resulted in the current way--clusters, Beowulf, and MPI. Now with limited gain in clock speed, we all wonder at how to harness the vast parallelism required to maintain the "capability computing" edge.

2022: Supercomputing Oddities

Satoshi Matsuoka (Tokyo Institute of Technology)

Bio: Satoshi Matsuoka has been a Full Professor at the Global Scientific Information and Computing Center (GSIC), a Japanese national supercomputing center hosted by the Tokyo Institute of Technology, since 2001. He received his Ph. D. from the University of Tokyo in 1993. He is the leader of the TSUBAME series of supercomputers, including TSUBAME2.0 which was the first supercomputer in Japan to exceed Petaflop performance and became the 4th fastest in the world on the Top500 in Nov. 2010, as well as the recent TSUBAME-KFC becoming #1 in the world for power efficiency for both the Green 500 and Green Graph 500 lists in Nov. 2013. He is also currently leading several major supercomputing research projects, such as the MEXT Green Supercomputing, JSPS Billion-Scale Supercomputer Resilience, as well as the JST-CREST Extreme Big Data. He has written over 500 articles according to Google Scholar, and chaired numerous ACM/IEEE conferences, most recently the overall Technical Program Chair at the ACM/IEEE Supercomputing Conference (SC13) in 2013. He is a fellow of the ACM and European ISC, and has won many awards, including the JSPS Prize from the Japan Society for Promotion of Science in 2006, awarded by his Highness Prince Akishino, the ACM Gordon Bell Prize in 2011, the Commendation for Science and Technology by the Minister of Education, Culture, Sports, Science and Technology in 2012, and recently the 2014 IEEE-CS Sidney Fernbach Memorial Award, the highest prestige in the field of HPC.

Abstract: Throughout my career my research topics were often not straightforward, but rather embodied various unconventional elements that stemmed from supposedly a variety of influences, such as my early years as a Commodore/Nintendo hacker, training in theoretical computer science and computer graphics/user interfaces, as well as curiosity for new but often wacky technologies. These and other elements blended together perhaps gave birth to early work which were at the time may have seemed unconventional but today have seen reincarnations in their modern forms. For example in the early 1990s I saw the potential of commodity 3D graphics

engines evolving into mainstream HPC computing elements, but discounted by most of the HPC field at the time, and it was true that most of the nascent work in the area was promising but not useful in practice. However as soon as programmable rendering pipelines in GPGPUs emerged the tide turned and it thankfully lead to TSUBAME series of supercomputers and this accolade whose credit should be shared with my collaborators and sponsors over the years. Now circa 2014 we are facing the exascale challenge and beyond, as well as extreme demands on big data leading to convergence of HPC and the latter infrastructures; these challenges are taxing the physical limits of our conventional IT, and likely not be able to overcome with traditional scale outs or evolutions in architectures and software. Rather, leading up to 2022 and beyond, the challenge is to be “odd”, or how to harness the unconventional so that we could overcome the challenge so that we can sustain our path to or even leapfrog exascale and bring on revolutionary changes to HPC, clouds, big data, etc. Such is where investments should be made and makes our involvement in supercomputing exciting.

What the \$#@! Is Parallelism? (And Why Should Anyone Care?)

Charles E. Leiserson (MIT)

Bio: Charles E. Leiserson received a B.S. from Yale University in 1975 and a Ph.D. from Carnegie Mellon University in 1981. He joined the MIT faculty in 1981, where he is now the Edwin Sibley Webster Professor of Electrical Engineering and Computer Science in the MIT Department of Electrical Engineering and Computer Science and head of the Supertech research group in the MIT Computer Science and Artificial Intelligence Laboratory.

Professor Leiserson’s research centers on the theory of parallel computing, especially as it relates to engineering reality. His Ph.D. thesis Area-Efficient VLSI Computation won the first ACM Doctoral Dissertation Award, as well as the Fannie and John Hertz Foundation’s Doctoral Thesis Award. He coauthored the first paper on systolic architectures. He invented the retiming method of digital-circuit optimization and developed the algorithmic theory behind it. On leave from MIT at Thinking Machines Corporation, he designed and led the implementation of the network architecture for the Connection Machine Model CM-5 Supercomputer, which incorporated the fat-tree interconnection network he developed at MIT. Fat-trees are now the preferred interconnect strategy for Infiniband technology. He introduced the notion of cache-oblivious algorithms, which exploit the memory hierarchy near optimally while containing no tuning parameters for cache size or cache-line length. He developed the Cilk multithreaded programming technology, which featured the first provably efficient work-stealing scheduler. He led the development of several Cilk-based parallel chess-playing programs, winning numer-

ous prizes in international competition. On leave from MIT as Director of System Architecture at Akamai Technologies, he led the engineering team that developed a worldwide content-distribution network numbering over 20,000 servers. He founded Cilk Arts, Inc., which developed the Cilk++ multicore concurrency platform and was acquired by Intel Corporation in 2009. Intel now embeds this technology in their Cilk Plus multithreaded programming environment. (See website for more on Professor Leiserson’s work.)

Abstract: Many people bandy about the notion of “parallelism,” saying such things as, “This optimization makes my application more parallel,” with only a hazy intuition about what they’re actually saying. Others cite Amdahl’s Law, which provides bounds on speedup due to parallelism, but which does not actually quantify parallelism. In this talk, I’ll review a general and precise quantification of parallelism provided by theoretical computer science, which every computer scientist and parallel programmer should know. I’ll also discuss why the impending end of Moore’s Law — the economic and technological trend that the number of transistors per semiconductor chip doubles every two years — will bring new poignancy to such theoretical concepts.

ACM Gordon Bell Finalists II

Chair: Subhash Saini (NASA Ames Research Center)

1:30pm-3pm

Room: New Orleans Theater

Anton 2: Raising the Bar for Performance and Programmability in a Special-Purpose Molecular Dynamics Supercomputer

David E. Shaw, J.P. Grossman, Joseph A. Bank, Brannon Batson, J. Adam Butts, Jack C. Chao, Martin M. Deneroff, Ron O. Dror, Amos Even, Christopher H. Fenton, Anthony Forte, Joseph Gagliardo, Gennette Gill, Brian Greskamp, C. Richard Ho, Douglas J. Ierardi, Lev Iserovich, Jeffrey S. Kuskin, Richard H. Larson, Timothy Layman, Li-Siang Lee, Adam K. Lerer, Chester Li, Daniel Killebrew, Kenneth M. Mackenzie, Shark Yeuk-Hai Mok, Mark A. Moraes, Rolf Mueller, Lawrence J. Nociolo, Jon L. Peticolas, Terry Quan, Daniel Ramot, John K. Salmon, Daniele P. Scarpazza, U. Ben Schafer, Naseer Siddique, Christopher W. Snyder, Jochen Spengler, Ping Tak Peter Tang, Michael Theobald, Horia Toma, Brian Towles, Benjamin Vitale, Stanley C. Wang, Cliff Young (D. E. Shaw Research)

Anton 2 is a second-generation special-purpose supercomputer for molecular dynamics simulations that achieves significant gains in performance, programmability, and capacity compared to its predecessor, Anton 1. The architecture of Anton 2 is tailored for fine-grained event-driven operation, which improves performance by increasing the overlap of computation with communication, and also allows a wider range of algorithms to

run efficiently, enabling many new software-based optimizations. A 512-node Anton 2 machine, currently in operation, is up to ten times faster than Anton 1 with the same number of nodes, greatly expanding the reach of all-atom biomolecular simulations. Anton 2 is the first platform to achieve simulation rates of multiple microseconds of physical time per day for systems with millions of atoms. Demonstrating strong scaling, the machine simulates a standard 23,558-atom benchmark system at a rate of 85 μ s/day—180 times faster than any commodity hardware platform or general-purpose supercomputer.

Award: Best Paper Finalist

24.77 Pflops on a Gravitational Tree-Code to Simulate the Milky Way Galaxy with 18600 GPUs

Jeroen Bédorf (Leiden Observatory, Leiden University), Evghenii Gaburov (SURFsara Amsterdam), Michiko S. Fujii (National Astronomical Observatory of Japan), Keigo Nitadori (RIKEN Advanced Institute for Computational Science), Tomoaki Ishiyama (Center for Computational Sciences, University of Tsukuba), Simon Portegies Zwart (Leiden Observatory, Leiden University)

We have simulated, for the first time, the long term evolution of the Milky Way Galaxy using 51 billion particles on the Swiss Piz Daint supercomputer with our N-body gravitational tree-code Bonsai. Herein, we describe the scientific motivation and numerical algorithms. The Milky Way model was simulated for 6 billion years, during which the bar structure and spiral arms were fully formed. This improves upon previous simulations by using 1000 times more particles, and provides a wealth of new data that can be directly compared with observations. We also report the scalability on both the Swiss Piz Daint and the US ORNL Titan. On Piz Daint the parallel efficiency of Bonsai was above 95%. The highest performance was achieved with a 242 billion particle Milky Way model using 18600 GPUs on Titan, thereby reaching a sustained GPU and application performance of 33.49 Pflops and 24.77 Pflops, respectively.

Thursday, November 20

Conference Awards Presentations

12:30pm - 1:30pm

Room: New Orleans Theater

Presenters: Barbara Chapman (University of Houston), John Grosh (Lawrence Livermore National Laboratory)

The SC14 conference awards, as well as selected ACM and IEEE awards, will be presented. Major accomplishments in supercomputing and related fields, from the most influential long-term research to outstanding results by young investigators, will be recognized by awards including the ACM Gordon Bell Prize, the SC14 Test of Time and the IEEE Technical Committee on Scalable Computing (TCSC) Young Investigators in Scalable Computing Awards; the SC14 Best Paper, Best Student Paper, and Best Poster Awards; George Michael Memorial HPC Ph.D. Fellowship; ACM Student Research Competition; and the Student Cluster Competition awards. This year, the recipient of the inaugural Best Visualization Award will also be recognized.

Birds of a Feather

Birds of a Feather

Don't just observe, ENGAGE! The Birds-of-a-Feather (BoF) sessions are among the most interactive, popular, and well-attended sessions of the SC Conference Series. BoFs provide a non-commercial, dynamic venue for conference attendees to openly discuss current topics of focused mutual interest within the HPC community with a strong emphasis on audience-driven discussion, professional networking and grassroots participation. SC14 continues this tradition with a full schedule of exciting BoFs.

We received 142 BoF proposals and a 22-person committee helped to select 85 for presentation. The selected BoFs span a wide range of timely topics. Technical topics include: accelerators, power/energy concerns, OS/Runtime issues and exascale application programming models. There will be several sessions on various HPC Center management topics including power and energy management, configuration tools and liquid cooling strategies. Lastly, the program includes a variety of non-technical BoFs on topics like research funding initiatives, the HPC workforce and even a behind the scenes look at what it takes to run SC each year.

BoF sessions will be held during the lunch hour and evening, Tuesday, November 18 through Thursday November 20.
(There are no Thursday evening BoFs.)

Birds of a Feather

Tuesday, November 18

Asynchronous Many-Task Programming Models for Next Generation Platforms

12:15pm-1:15pm

Room: 396

Robert Clay, Janine Bennett (Sandia National Laboratories)

Next generation platform architectures will require us to fundamentally rethink our programming and execution models due to a combination of factors including extreme parallelism, data locality issues, and resilience. The asynchronous, many-task (AMT) programming and execution model is emerging as a leading new paradigm, with many variants of this new model being proposed. This BoF will bring together a panel of experts to survey the start-of-the-art in AMT runtime systems. We invite researchers, users, and developers to participate in this session to form collaborations, learn about, and influence the direction of this body of work.

Chapel Lightning Talks 2014

12:15pm-1:15pm

Room: 293

Sung-Eun Choi (Cray, Inc.), Richard Barrett (Sandia National Laboratories)

Are you a scientist considering a modern high-level language for your research? Are you a language enthusiast who wants to stay on top of new developments? Are you an educator considering using Chapel in your classroom? Are you already a Chapel fan and wondering what's in store for the future? Then this is the BOF for you!

In this BOF, we will hear "lightning talks" on community activities involving Chapel. We will begin with a talk on the state of the Chapel project, followed by a series of talks from the broad Chapel community, wrapping up with Q&A and discussion.

Code Optimization War Stories

12:15pm-1:15pm

Room: 298-99

Georg Hager (Erlangen Regional Computing Center), Gerhard Wellein (University of Erlangen), Jan Treibig (Erlangen Regional Computing Center)

Code optimization is essential for the efficient use of highly parallel systems, small clusters, and even desktop machines. Although there is a trend in HPC research towards automatic tuning tools, most of these efforts are still hard work done

by domain scientists or performance experts in computing centers. However, code optimization is often not regarded as novel scientific contribution, so many of these activities go unpublished and unnoticed beyond a very small community. With this BoF we want to reach developers from diverse fields of domain and computer science to share and discuss interesting, educational "war stories" in code optimization.

Experimental Infrastructures for Open Cloud Research

12:15pm-1:15pm

Room: 388-89-90

Kate Keahey (Argonne National Laboratory), Dan Stanzione, Warren Smith (University of Texas at Austin)

Cloud Computing has emerged as a critical infrastructure for scientific, enterprise, and commercial computing. To support the creation of such infrastructure, there is a need for quality testbeds for development and testing. Commercial companies create their own testbeds, but academic and government cloud researchers don't have access to them. Therefore, funding agencies around the world have or soon will construct an array of new experiment infrastructures for in cloud computing. This BOF will provide a forum for the providers of these experimental infrastructures and their potential users to come together, and form a community of testbeds.

HDF5: State of the Union

12:15pm-1:15pm

Room: 398-99

Quincey Koziol, Mohamad Chaarawi (HDF Group)

This BoF is a forum for HDF5 developers and users to interact. HDF5 developers will describe the current status of HDF5 and discuss future plans, followed by an open discussion.

HPC System and Data Center Energy Efficiency Metrics and Workloads

12:15pm-1:15pm

Room: 294

Daniel Hackenberg (Technical University Dresden), Robin Goldstone (Lawrence Livermore National Laboratory), Nicolas Dubé (Hewlett-Packard)

US House Bill 540 just passed requiring energy efficiency metrics for federal data centers. Without heavy lifting, the metric will be PUE. However, with PUEs approaching 1.1, PUE may have limited utility.

Should and can we move on from PUE? This BoF will discuss what's needed for driving the next level of improvement. It will also discuss how to advance system-level energy efficiency workloads and metrics from the currently-used HPL FLOPS/Watt.

The BoF will host a panel and field questions from both a moderator and the floor. Chung-Hsing Hsu from ORNL will moderate, and panelists will include experts in the field.

LLVM in HPC: Uses and Desires

12:15pm-1:15pm

Room: 283-84-85

Hal Finkel (Argonne National Laboratory), Jim Cownie (Intel Corporation)

The LLVM compiler infrastructure is widely used for software tool development from compilers through source to source translators, dynamic language implementations, and even in an MPI runtime.

An elite group of LLVM users will describe their use of LLVM (in three slides and three minutes each), and the audience will then be able to question them, and suggest what features they would like in LLVM and its associated tools.

Lustre Community BoF: At the Heart of HPC and Big Data

12:15pm-1:15pm

Room: 275-76-77

Galen Shipman (Oak Ridge National Laboratory), Hugo Falter (European Open File System)

Lustre is the leading open source file system for HPC. Lustre is the only open-source, community developed file system in the top 10 of the top 500 HPC systems. Lustre is now at the heart of many HPC and Big Data infrastructures supporting multiple sectors from financial services, oil and gas, advanced manufacturing, advanced web services, to basic science. At this year's Lustre Community BOF the worldwide community of Lustre developers, administrators, and solution providers will gather to discuss new challenges and corresponding opportunities emerging within HPC and Big Data and how Lustre can continue to evolve to support them.

Ninth Graph500 List

12:15pm-1:15pm

Room: 286-87

Richard Murphy (Micron Technology, Inc.), David Bader (Georgia Institute of Technology), Andrew Lumsdaine (Indiana University)

Large-scale data analytics applications represent increasingly important workloads but most of today's supercomputers are ill suited to them. Backed by a steering committee of over

30 international HPC experts from academia, industry, and national laboratories, Graph500 works to establish and administer large-scale benchmarks that are representative of these workloads. This BOF will unveil the ninth Graph500 list and the forth Green Graph500 list.

Operating System and Run-Time for Exascale

12:15pm-1:15pm

Room: 391-92

Marc Snir (Argonne National Laboratory), Arthur Maccabe (Oak Ridge National Laboratory), John Kubiawicz (University of California, Berkeley)

DOE's Advanced Scientific Computing Research (ASCR) program funded in 2013 three major research projects on Operating System and Runtime (OS/R) for extreme-scale computing: ARGO, HOBBS, and X-ARCC. These three projects share a common vision of the OS/R architecture needed for extreme-scale and explore different components of this architecture, or different implementation mechanisms.

The goals of the BOF are to engage researchers and vendors in this effort and obtain their feedback. We shall present the common vision and current status of the projects and discuss mechanisms for collaboration, standardization and technology transfer.

QuantumChemistry500

12:15pm-1:15pm

Room: 295

Maho Nakata (RIKEN), Toshiyuki Hirano (University Of Tokyo), Kazuya Ishimura (Institute For Molecular Science)

Quantum chemistry or computational chemistry is now widely used by researchers and many good program packages are available. However, there are still many issues for massively parallel implementations. In this BoF, we would like to establish a performance index for quantum chemistry programs which we will discuss and consider how fast SCF (or other methods) calculations are, and how large molecules we can calculate, and most importantly, how parallel program packages are, and which systems can calculate large molecules fastest?

Resilience and Power-Efficiency Challenges at Exascale: What Benchmarking, System Monitoring Tools and Hardware Support Do We Need?

12:15pm-1:15pm

Room: 292

Devesh Tiwari, Saurabh Gupta (Oak Ridge National Laboratory)

As we approach exascale, resilience and power efficiency are anticipated to be one of the most critical challenges to overcome. While we have made significant advances in terms

of theoretical research, development and adoption of practical tools in this area are still in their adolescence. Therefore, the goal of this BoF is to brainstorm about what benchmarking and system monitoring tools we need to develop and deploy as a community (from HPC application developers to HPC system operation teams). Outcomes from this BoF, desired hardware or system support, will be forwarded to vendors to influence the technology road-map of vendors.

Scalable Adaptive Graphics Environment (SAGE) for Global Collaboration

12:15pm-1:15pm

Room: 393-94-95

Jason Leigh (University of Hawaii at Manoa), Maxine Brown, Luc Renambot (University of Illinois at Chicago)

SAGE, the Scalable Adaptive Graphics Environment, is the de facto operating system for managing Big Data content on scalable-resolution tiled display walls, providing the scientific community with persistent visualization and collaboration services for global cyberinfrastructure. SAGE2, the next-generation SAGE, is being introduced to the SAGE global user community and to potential users at the SC14 BOF; SAGE2 is browser-based and integrates cloud services, a multi-resolution image viewer, a stereoscopic image viewer, mirrored portals, and network streaming of high-resolution imagery. SC provides an unparalleled community-building opportunity for developers and users to meet, share use cases, and discuss user requirements and future roadmaps.

SIGHPC Annual Meeting

12:15pm-1:15pm

Room: 383-384-385

The 2014 HPC Challenge Awards

12:15pm-1:15pm

Room: 273

Piotr Luszczek (University of Tennessee, Knoxville), Jeremy Kepner (Massachusetts Institute of Technology Lincoln Laboratory)

The 2014 HPC Challenge Awards BOF is the 10th edition of an award ceremony that seeks high performance results in broad categories taken from the HPC Challenge benchmark as well as elegance and efficiency of parallel programming and execution environments. The performance results come from the HPCC public database of submitted results that are unveiled at the time of BOF. The competition for the most productive (elegant and efficient) code takes place during the BOF and is judged on the spot with winners revealed at the very end of the BOF. Judging and competing activities are interleaved to save time.

Towards European Leadership in HPC

12:15pm-1:15pm

Room: 386-87

Jean-François Lavignon (Bull), Jean-Philippe Nominé (French Alternative Energies and Atomic Energy Commission), Marcin Ostasz (Barcelona Supercomputing Center)

The development of European HPC is gaining momentum. With the creation of the contractual Public-Private Partnership (cPPP) for HPC, Europe has now laid foundations for a comprehensive program that will result in a 700M Euros investment from the European Commission in this technology within the Horizon 2020 Research and Development Programme.

The objective of this BOF is to present the European HPC Ecosystem in this context: the technology supply chain, the infrastructure and the application resources and its development plans with a view to compare them with those of other regions and facilitate international collaborations.

Compiler Vectorization and Achieving Effective SIMD for HPC Software

5:30pm-7pm

Room: 391-92

David Mackay, Xinmin Tian (Intel Corporation)

SIMD operations abound in high performance computing. When utilized effectively SIMD delivers both high performance and power efficient computing. Matching the SIMD operations to the hardware SIMD execution units is still challenging for many HPC software packages. How can developers express code so compilers or other runtime environments effectively exploit SIMD operations is just as important today as it was many years ago when vector computers were first developed. This BOF will have a panel of compiler developers to discuss the state of compilers as well as software developers to express user perspective.

Design, Commissioning and Controls for Liquid Cooling Infrastructure

5:30pm-7pm

Room: 383-84-85

David Martinez (Sandia National Laboratories), Marriann Silveira (Lawrence Livermore National Laboratory), Herbert Huber (Leibniz Supercomputing Center)

Liquid cooling is key to dealing with heat density, energy efficiency, and increasing the performance of supercomputers. It becomes even more predominant looking forward. The transition to liquid cooling, however, comes with challenges.

Each system and each site comes with its specific issues regarding liquid cooling. These can complicate procurements, design, installation, operations, and maintenance.

This BoF will review the immediate past history from a lessons learned perspective as well as discuss what's needed for liquid cooling to be implemented more readily in the future. It will explore HPC center's liquid cooling design parameters, commissioning guidelines and control systems.

Forming a Support Organization for HPC User Service Providers

5:30pm-7pm

Room: 297

James Lupo (Louisiana State University), Dana Brunson (Oklahoma State University), Galen Collier (Clemson University)

Advanced cyberinfrastructure for academia, government and industry involves high performance computing (HPC) and necessitates delivery of resources and services to support diverse community needs. This BOF will focus discussion on the foundational elements for mutual aid which cross organizational boundaries. We invite everyone interested in improving HPC support services to join this discussion. Identifying ideas for organizational scale, funding strategies, metrics of success, and ROI for stakeholders are some of the desired outcomes. The XSEDE Campus Champions and the Advanced Cyberinfrastructure Research and Education Facilitators projects will make a series of short presentations, followed by an open general discussion.

From Big Data to the Long Tail: Publishing Computational Data

5:30pm-7pm

Room: 286-87

Raymond Plante (University of Illinois at Urbana-Champaign), Beth Plale (Indiana University), Robert Pennington (University of Illinois at Urbana-Champaign)

The ability to re-use data and validate research results have improved dramatically with the rise of the Internet as a vital tool of modern research. It is now possible to publish digital data alongside of the scientific results they support. As many disciplines have begun to make significant strides data sharing and publishing, it is becoming clear that there are benefits to be gained from a cross-disciplinary approach to data publishing. In this BOF, we will discuss what is needed to enable the publishing of computational science data and how those needs might be addressed in a cross-disciplinary framework.

Here come OpenMP 4 Implementations and OpenMPCon

5:30pm-7pm

Room: 291

Michael Wong (IBM Corporation), Bronis de Supinski (Lawrence Livermore National Laboratory)

In July of last year, the OpenMP Architecture Review Board released version 4.0 of the specification of the OpenMP API. Since then, implementers have been hard at work developing compilers for this API. The current implementations of OpenMP 4.0 will be presented to users, so that they become aware of the great opportunities that OpenMP 4.0 can give to them. Plans for an OpenMPCon and for future standard extensions will be discussed with the audience. The intended audience consists of users and implementers. The format consists of presentations by users and implementers, discussions, and a lively crossfire session with implementers.

How TORQUE is Changing to Meet the New Demands for HPC

5:30pm-7pm

Room: 293

Kenneth Nielson, David Beer, Jill King (Adaptive Computing)

TORQUE has been actively improving to better meet the challenges of an ever evolving HPC world. We will be introducing new features added since last year, such as power management, task placement and resource enforcement as well as showing new improvements in performance and scaling.

As with all TORQUE Birds of a Feather, we will be engaging the TORQUE community for feedback and ideas on what will be needed to keep TORQUE at the forefront of HPC management.

HPC Systems Engineering, Suffering, and Administration

5:30pm-7pm

Room: 294

William Scullin (Argonne National Laboratory), Jenett Tillotson (Indiana University), Adam Hough (Petroleum Geo-Services)

Systems chasing exascale often leave their administrators chasing yottascale problems. This BOF is a forum for the administrators, systems programmers, and support staff behind some of the largest machines in the world to share solutions and approaches to some of their most vexing issues and meet other members of the community. This year we are focusing on how to best share with others in the community; monitoring tools and accounting; training users, new administrators, and peers; management of expectations; and discussing new tools, tricks, and troubles.

Integrated Optimization of Performance, Power and Resilience for Extreme Scale Systems**5:30pm-7pm****Room: 298-99***Dong Li (Oak Ridge National Laboratory), Kathryn Mohror (Lawrence Livermore National Laboratory)*

Performance, power and resilience are three key system design factors for extreme scale systems. However, they are often considered in isolation, and we are largely blind to the integrated impact of the three factors for future HPC systems. Hence, it is urgent for experts in the HPC community to jointly explore the interactions of these three factors. This BOF will bring together researchers from the fields of performance optimization, power-aware computing, and fault tolerance. We will discuss timely topics, including cross-layer system design, the impact of future hardware, and modeling methods and evaluation metrics to achieve the integrated system optimization.

Joint NSF/ENG and AFOSR Initiative on Dynamic Data Systems**5:30pm-7pm****Room: 283-84-85***Frederica Darema (Air Force Office of Scientific Research), Zhi Tian (National Science Foundation)*

This BOF presents the Joint NSF/ENG and AFOSR Initiative to support multidisciplinary research aimed to transform our ability to understand, manage and control the operation of complex, multi-entity natural or engineered systems, through innovative approaches that consider new dimensions in Big Data and in Big Computing, and a symbiotic combination of Data and Computing. Big Computing is the integrated set of platforms spanning the high-end/mid-range computing and computing on multitudes of sensors and controllers, holistically viewed as a unified platform, and Big Data encompasses HPC data and real-time and archival data relating to ubiquitous sensing and control in data-intensive systems.

Migration to Lustre 2.5 and Beyond: Best Practices & Shared Experiences**5:30pm-7pm****Room: 292***Sharan Kalwani, Amitoj Singh (Fermi National Accelerator Laboratory)*

Capture Migration experiences going to Lustre 2.5. Intended for all sites using Lustre--all versions---and planning to or had made the leap forward to version 2.5 and beyond. Shared experiences are key, including planning exercises, how to deal with unexpected situations and put together best practices.

MPICH: A High-Performance Open-Source MPI Implementation**5:30pm-7pm****Room: 386-87***Ken Raffenetti, Pavan Balaji, Rajeev Thakur (Argonne National Laboratory)*

MPICH is a widely used, open-source implementation of the MPI message passing standard. It has been ported to many platforms and used by several vendors and research groups as the basis for their own MPI implementations. This BoF session will provide a forum for users of MPICH as well as developers of MPI implementations derived from MPICH to discuss experiences and issues in using and porting MPICH. Future plans for MPICH will be discussed. Representatives from MPICH-derived implementations will provide brief updates on the status of their efforts. MPICH developers will also be present for an open forum discussion.

OpenCL: Version 2.0 and Beyond**5:30pm-7pm****Room: 275-76-77***Tim Mattson (Intel Corporation), Simon McIntosh-Smith (University of Bristol)*

OpenCL enables heterogeneous computing (e.g. a combination of CPUs, GPUs, and coprocessors) without locking users into a single vendor's products. In this BOF, we will discuss the latest developments in OpenCL including the recent OpenCL 2.0 specification. We will then convene a panel representing divergent views on the future of OpenCL to engage in an interactive debate on where to take OpenCL. We will close the BOF with a survey of the audience to create "hard data" we can present to the OpenCL standards committee to help them address the unique needs of the HPC community.

OpenPower: Future Directions with HPC**5:30pm-7pm****Room: 273***Scot Schultz (Mellanox Technologies), Dirk Pleiter (Juelich Research Center)*

OpenPOWER Foundation was founded in 2013 as an open technical membership organization to allow greater collective innovation around the entire ecosystem of POWER-based CPUs. Active collaboration activities have been launched both on hardware and software, in particular in the arena of tightly coupled heterogeneous compute nodes. Initial OpenPOWER designs have become available and are presented elsewhere at this conference.

We will present the Foundation, introduce established working groups, discuss roadmaps and goals with conference participants and invite involvement. We will also focus how the foundation can better include the HPC community via targeted OpenPOWER workgroups and initiatives.

Performance Analysis and Simulation of MPI Applications and Runtimes at Exascale

5:30pm-7pm

Room: 398-99

Frederic Suter (National Institute of Nuclear and Particle Physics, French National Center for Scientific Research)

Scaling MPI applications and runtimes to exascale and hierarchies of millions of heterogeneous cores is a true challenge. It requires performance analysis made even more complex by the lack of (projections of) such exascale systems. However, many tools have been developed to this end, but in an uncoordinated way, each having its own focus, strengths, and limitations.

This BoF has two goals. First, we invite potential users for an informative overview of the existing tools through a series of short talks. Second, we gather tool developers to set up a working group to favor interoperability and foster cross-fertilization.

PRObE - 1500 New Machines for Systems Research

5:30pm-7pm

Room: 396

Andree Jacobson (New Mexico Consortium), Garth Gibson (Carnegie Mellon University)

PRObE - The NSF sponsored Parallel Reconfigurable Observational Environment - provides a unique, large scale, systems research facility for the HPC community. A current technology refresh brought us 1500 new nodes currently being integrated into the testbed and gradually becoming available for use by the community - free of charge. If your systems related project require a large scale testbed and you need dedicated access to hardware - remotely, or in person - this forum is for you! We will show the current state of our systems and perform a mini-howto demonstration for how to use our clusters.

Reconfigurable Supercomputing

5:30pm-7pm

Room: 393-94-95

Martin Herbordt (Boston University), Alan George (University of Florida), Herman Lam (University of Florida)

Reconfigurable supercomputing (RSC) is characterized by hardware that adapts to match the needs of each application, offering unique advantages in speed per unit energy. With a proven capability of 2 PetaOPS at 12KW, RSC has an important role to play in the future of high-end computing. The Novo-G in the NSF CHREC Center at the University of Florida is rapidly moving towards production utilization in various scientific and engineering domains. This BoF introduces the architecture of such systems, describes applications and tools being developed, and provides a forum for discussing emerging opportunities and issues for performance, productivity, and sustainability.

Strategies for Academic HPC Centers

5:30pm-7pm

Room: 388-89-90

Jackie Milhans, Joseph Paris (Northwestern University), Sharon Broude Geva (University of Michigan)

The purpose of this discussion is to explore challenges of smaller, academic HPC centers. There are three general topics of discussion: funding and sustainability, community outreach and user support, and hiring, recruitment and retention. This BoF will showcase team presentations from 5 universities on successes, current struggles, and future plans. After each presentation, the BoF will be open for comments and questions; the second half of the BoF will be used to continue the discussion. The intended audience for this BOF is university staff involved in research computing (primarily HPC), such as directors and managers, support staff, and students.

TOP500 Supercomputers

5:30pm-7pm

Room: New Orleans Theater

Erich Strohmaier (Lawrence Berkeley National Laboratory)

The TOP500 list of supercomputers serves as a “Who’s Who” in the field of HPC. It started as a list of the most powerful supercomputers in the world and has evolved to a major source of information about trends in HPC.

This BoF will present detailed analyses of the TOP500 and discuss the changes in the HPC marketplace during the past years. The BoF is meant as an open forum for discussion and feedback between the TOP500 authors and the user community.

Women in HPC: Mentorship and Leadership

5:30pm-7pm

Room: 295

Rebecca Hartman-Baker (iVEC), Fernanda Foertter (Oak Ridge National Laboratory), Toni Collis (EPCC)

Although women comprise more than half the world’s population, they make up only a small percentage of people in the Science, Technology, Engineering, and Mathematics (STEM) fields, including the many disciplines associated with high-performance computing (HPC). Studies have shown that many minorities in STEM attribute their success to important role models such as teachers and mentors within the field. Despite women’s low participation in HPC, there are a number of notable women in the field. In this BOF, we bring successful women leaders in HPC to discuss their career paths and the role of mentorship in their success.

Wednesday, November 19

Application Readiness and Portability for Leadership Computing

12:15pm-1:15pm

Room: 293

*Tjerk Straatsma (Oak Ridge National Laboratory),
Timothy Williams (Argonne National Laboratory)*

This Birds-of-a-Feather session presents the application development partnerships the Oak Ridge and Argonne Leadership Computing Facilities propose for the next round of application readiness programs for supercomputers expected at their sites in 2017-2018. Through open calls for proposals, ideas will be solicited for pre-production projects with access to the future machines to conduct early high-impact science. Projects will focus on application development, performance portability and programming strategies to enable that science. The BoF will engage leaders in computational science in discussing the challenges in achieving scalability and performance portability of scientific applications for the next pre-exascale computers.

Dynamic Power Management for MW-Sized Supercomputer Centers

12:15pm-1:15pm

Room: 292

*Ghaleb Abdulla (Lawrence Livermore National Laboratory),
Terry Hewitt (Science and Technology Facility Center),
Herbert Huber (Leibniz Supercomputing Center)*

Supercomputing centers are poised to begin a transition to “dynamic power management.” Multiple forces are driving this need with energy cost control in the forefront. Supercomputer systems have increasingly rapid, unpredictable and large power fluctuations with costly implications for the supercomputer center. In addition, electricity service providers may incentivize supercomputing centers to change their timing and/or magnitude of demand and to help address variable renewable generation. To adapt to this new landscape, supercomputing centers may employ strategies to dynamically and in real-time control their electricity demand. This BoF seeks those interested in sharing experiences with dynamic power and energy management.

Getting Scientific Software Installed: Tools & Best Practices

12:15pm-1:15pm

Room: 298-99

Andy Georges, Jens Timmerman, Ewan Higgs (Ghent University)

We intend to provide a platform for presenting and discussing tools to cope with the ubiquitous problems that come forward when building and installing scientific software, which is known to be a tedious and time consuming task.

Several user support tools for allowing scientific software to be installed and used will briefly be presented, for example (but not limited to) EasyBuild (UGent), Lmod (TACC) and Spack (LLNL).

We would like to bring various experienced members of HPC user support teams and system administrators as well as users together for an open discussion on tools and best practices.

High Performance Geometric Multigrid (HPGMG) Proposal for a New Top500 Benchmark

12:15pm-1:15pm

Room: 294

*Mark Adams (Lawrence Berkeley National Laboratory),
Jed Brown (Argonne National Laboratory), Sam Williams
(Lawrence Berkeley National Laboratory)*

This is the first SC BOF for HPGMG, a proposed a Top500 benchmark. The goal of this BOF is to engage the community by describing the project, presenting internal and public experience with HPGMG, and eliciting community input with open discussion. We will briefly describe the proposal and the current state of the project, present experience from the HPGMG team and the public (interested parties are encouraged to contact us about presenting), and provide a forum to elicit public comment and discussion on the HPGMG effort. Community members interested in participating can contact us at hpgmg.org.

HPC Debugging Techniques

12:15pm-1:15pm

Room: 393-94-95

*Chris Gottbrath (Rogue Wave Software), Dong Ahn
(Lawrence Livermore National Laboratory), Mike Ashworth
(STFC Daresbury Laboratory)*

This BOF will bring together users interested in debugging parallel programs on HPC systems for the purposes of sharing experiences and identifying and disseminating best practices. HPC programs are much more complex to debug than the typical serial application: they run in batch environments, in run-time environments which are optimized to deliver peak performance, are often memory constrained and contain massive parallelism in the forms of processes, threads, accelerators and vectors.

Come to this BOF ready to learn best practices and share your experiences. We'll have a few brief talks and then an open discussion.

HPC Education: Meeting of the New SIGHPC Education Chapter

12:15pm-1:15pm

Room: 295

Steven Gordon (Ohio Supercomputer Center), David Halstead (National Radio Astronomy Observatory)

The newly formed SIGHPC Education chapter has as its purpose the promotion of interest in and knowledge of applications of High Performance Computing (HPC). This BOF will bring together those interested in promoting the education programs through the formal and informal activities of the chapter. The current officers will present information on the chapter organization, membership, and an initial set of proposed activities. They will then lead an open discussion from participants to solicit their ideas and feedback on chapter activities.

HPC Job Scheduling Challenges and Successes from the PBS Community

12:15pm-1:15pm

Room: 291

Bill Nitzberg (Altair), Greg Matthews (Computer Sciences Corporation / NASA Ames Research Center)

More than 20 years ago, NASA developed the PBS software. Now, PBS is consistently among the top 3 most-reported technologies for HPC job scheduling (as reported by IDC). But scheduling is hard -- every site has unique processes, unique goals, and unique requirements.

Join fellow members of the PBS community to share challenges and solutions and learn others' tips and tricks for HPC scheduling in the modern world, including exascale, clouds, power management, GPUs, Xeon Phi, Big Data, and more.

Open MPI State of the Union

12:15pm-1:15pm

Room: 283-84-85

Jeffrey Squyres (Cisco Systems), George Bosilca (University of Tennessee)

It's been another great year for Open MPI. We've added new features, improved performance, and completed our journey towards a full MPI-3.0 implementation. We'll discuss what Open MPI has accomplished over the past year and present a roadmap for the next year.

One of Open MPI's strengths lies in its diverse community: we represent many different viewpoints from across the spectrum of HPC users. To that end, we'll have a Q&A session to address questions and comments from our community.

Join us at the BOF to hear a state of the Union for Open MPI. New contributors are welcome!

PGAS- The Partitioned Global Address Space Programming Model

12:15pm-1:15pm

Room: 273

Tarek El-Ghazawi (George Washington University), Lauren Smith (U.S. Government)

The partitioned global address space (PGAS) programming model strikes a balance between the ease of programming due to its global address model and efficiency due to locality awareness. In addition to scalable systems, PGAS can become even more widely acceptable as latency matters with non-uniform cache and memory access in manycore chips. There are many active efforts for specific PGAS paradigms such as Chapel, UPC and X10. The PGAS BoF at SC14 will bring researchers and practitioners from those efforts for cross-fertilization of new ideas, and to address common issues of concern and common infrastructures.

Power API for HPC: Standardizing Power Measurement and Control

12:15pm-1:15pm

Room: 286-87

Stephen Olivier, Ryan Grant, James Laros (Sandia National Laboratories)

The HPC community faces considerable constraints on power and energy of HPC installations going forward. A standardized, vendor-neutral API for power measurement and control is sorely needed for portable solutions to these issues at the various layers of the software stack. In this BOF, we discuss such an API, which has been designed by Sandia National Laboratories and reviewed by representatives of Intel, AMD, IBM, Cray, and other industry, laboratory, and academic partners. The BOF will briefly introduce the API and feature an interactive panel discussion with experts from among these partners, with ample time for audience questions and comments.

Programming Abstractions for Data Locality

12:15pm-1:15pm

Room: 391-92

Didem Unat (Koç University), John Shalf (Lawrence Berkeley National Laboratory), Torsten Hoefler (ETH Zürich)

This BoF discusses the abstractions that libraries, languages, and compilers can provide to move away from a compute-centric to a data-centric programming model as a way to address Exascale challenges. The attendees of the BoF previously participated in "Programming Abstractions for Data Locality" workshop in April 2014 (Switzerland) and compiled a report summarizing their findings (<http://www.padalworkshop.org>). The goal of the BoF is to share these findings with the rest of the community and highlight the key points. This BoF will foster collaboration across this community and help pave the way towards standardization of some of the programming concepts.

Python for High Performance and Scientific Computing**12:15pm-1:15pm****Room: 275-76-77**

*Andreas Schreiber (German Aerospace Center),
William Scullin (Argonne National Laboratory),
Andy Terrel (Continuum Analytics, Inc.)*

This BoF is intended to provide current and potential Python users and tool providers in the high performance and scientific computing communities a forum to talk about their current projects; ask questions of experts; explore methodologies; delve into issues with the language, modules, tools, and libraries; build community; and discuss the path forward.

SC Community Town Hall**12:15pm-1:15pm****Room: 396**

William Gropp (University of Illinois at Urbana-Champaign)

Have you wondered how the SC conference works? Do you have some ideas for improving the conference? Are you interested in becoming part of the SC organization? This session is a “town hall” meeting for the organizers and the attendees of SC to meet and discuss the conference. The session will start with a brief presentation about the SC organization and how the conference works. Following that will be an open discussion session. Members of the SC steering committee and chairs of the upcoming 2015 and 2016 SC Conferences will lead the discussions.

Scalable Industrial Software: Future or Fantasy?**12:15pm-1:15pm****Room: 386-87**

Fred Streitz (Lawrence Livermore National Laboratory), Rick Arthur (General Electric)

We invite independent software providers, hardware vendors, universities, industry end users, and members of the federal government to discuss opportunities for fostering innovation in software development and productivity for extreme computing systems for economic outcome.

The Open Science Data Cloud and PIRE Fellowships: Handling Scientific Datasets**12:15pm-1:15pm****Room: 297****The Open Science Data Cloud and PIRE Fellowships: Handling Scientific Datasets**

Robert Grossman, Maria Patterson (University of Chicago)

A wide variety of disciplines are producing unprecedented volumes of data, and research groups are struggling to manage, analyze, and share their medium to large size datasets. The Open Science Data Cloud (OSDC) is an open source infrastructure designed to allow scientists to easily manage, analyze, integrate and share medium to large size scientific datasets. OSDC is operated and managed by the not-for-profit Open Cloud Consortium. Come to this session to learn about how you can use the OSDC for your big data research projects and how to participate in the NSF sponsored OSDC PIRE (Partnerships in Research Education) program.

The OpenPOWER Foundation: Addressing HPC Challenges Through the POWER Eco-System**12:15pm-1:15pm****Room: 383-84-85**

*Stephen Poole (Oak Ridge National Laboratory),
Kathryn O'Brien (IBM), Duncan Poole (NVIDIA Corporation)*

In 2013, the OpenPOWER Foundation was formed to facilitate innovative system solutions based on the POWER architecture. Some of its members include IBM, NVIDIA, Mellanox, Altera, Google, Tyan, Micron, and Samsung. OpenPOWER has a broader mandate than HPC but in this BoF we focus on collaboration between the HPC community and the OpenPOWER effort to address HPC challenges related to system architecture, networking, memory designs, and programming models. Being a novel architecture suitable for HPC, we expect this session to be well attended and instrumental in building a HPC community for the OpenPOWER foundation.

Application Experiences with Emerging PGAS APIs: MPI-3, OpenSHMEM and GASPI**5:30pm-7pm****Room: 386-87**

Hans-Christian Hoppe, Marie-Christine Sawley (Intel Corporation), Christian Simmendinger (T-Systems)

PGAS languages and APIs have been around for some time, and PGAS advantages over message passing have been studied in depth. Yet, widespread use of PGAS by HPC developers has not happened. Advances in PGAS APIs promise to significantly increase PGAS use, avoiding the effort and risk involved in adopting a new language, and promising higher efficiency on modern systems than message passing can achieve.

This BOF gives a concise update on progress on PGAS communication APIs, presents recent experiences in porting applications to these interfaces, and discusses how PGAS can evolve to maximize take-up by application developers.

Challenges in Managing Small HPC Centers

5:30pm-7pm

Room: 295

Elizabeth Leake (University of Wisconsin-Milwaukee), Beth Anderson (Intel Corporation), Kevin Walsh (University of Wisconsin-Milwaukee)

This session provides a venue for those involved in managing small/campus-scale, HPC capacity clusters to share their challenges, woes and successes. Attendees will also react to the findings of a pre-conference survey that investigated the compute, storage, submit/compile, backup and scheduler environments that are common in systems of this size. Trend data from four consecutive years will be provided to catalyze discussion.

Chapel Users Group Meeting

5:30pm-7pm

Room: 383-84-85

Bradford Chamberlain, Sung-Eun Choi (Cray, Inc.)

Chapel is an emerging parallel programming language being developed as an open-source project with contributions from industry, academia and government labs. The Chapel Users Group (CHUG) meeting will begin with an overview of the language and current development status, followed by an interactive discussion and Q&A with the audience. Discussion topics may include growing the Chapel user community, teaching Chapel to both students and programmers, and creating a Chapel Foundation to oversee and govern the language and implementation.

Keeping in the spirit of previous CHUG meetings, afterwards we will proceed to a local pub for the annual CHUG happy hour.

Codesign for the Department of Energy's Computational Science Community

5:30pm-7pm

Room: 292

Richard Barrett (Sandia National Laboratories), Charles Still (Lawrence Livermore National Laboratory), Allen McPherson (Los Alamos National Laboratory)

The Department of Energy high performance computing community is actively engaged in codesign efforts as a means of ensuring that new architectures more effectively support mission critical workloads. The combinations of application

codes, represented by proxy programs, algorithms, programming models, system software, and architecture provides a concrete and powerful means for interacting with vendors. In this session we will present and discuss concrete examples of the impact of these efforts. We invite members of the computational science community to participate in this session, as a means of learning about, forming collaborations, and influencing the direction of this work.

Developing the HPC Workforce of the Future

5:30pm-7pm

Room: 393-94-95

Roscoe Giles (Boston University), Barbara Chapman (University of Houston), Scott Lathrop (Shodor / University of Illinois at Urbana-Champaign)

This BOF will stimulate a broader community discussion on the HPC workforce and share ideas on how it might be expanded in both size and quality.

We start with a brief summary of the recent DOE ASCAC report that analyzes the workforce challenges facing the DOE national labs in High Performance Computing and related fields in the Computing Sciences. Action beyond the scope of a single agency, program, or institution is needed to create the workforce of the future.

The rest of the BOF will be an open discussion among key stakeholders and interested parties from the SC community.

Evaluation Infrastructure for Future Large-Scale Chips and HPC Systems

5:30pm-7pm

Room: 391-392

George Micheliogiannakis, Farzad Fatollahi-Fard, Dave Donofrio (Lawrence Berkeley National Laboratory)

Future HPC nodes will likely contain processors with thousands of cores and will spark a revolution in virtually all areas of computing. Understanding the complex tradeoffs of this radical shift in computer architecture and its impact on applications will require advanced, multi-resolution simulation and emulation infrastructures. This BoF will focus on the challenges and fundamental tradeoffs in evaluating future many-core chips and exascale systems. The goal is to spark a lively discussion and exchange of opinions and experiences with 5-minute conversation starter presentations on a range of topics including simulation accuracy, fidelity and feasibility.

Experiences in Delivering a Campus Data Service Using Globus

5:30pm-7pm

Room: 298-99

Ian Foster (University of Chicago, Argonne National Laboratory), Steve Tuecke, Rachana Ananthakrishnan (University of Chicago)

Universities and other non-profit research institutions are exploring options for delivering secure, high-performance services to help researchers deal with the growing volume of data on their campuses. Globus is currently used for file transfer and sharing at over 100 institutions, and by national cyberinfrastructure providers such as XSEDE and NERSC. This BoF will convene campus computing and other HPC practitioners to discuss their experiences with Globus and related technologies for research data management. We will provide a brief update on Globus and invite select Globus users to present their use cases as context for an open discussion with the audience.

Integrating HPC Technologies in BigData Architectures

5:30pm-7pm

Room: 283-84-85

Ryan Quick, Arno Kolster (PayPal)

The second Integrating HPC Technologies (HPC BigData) BoF sharpens community focus through guided collaboration. Fellowship between industry, academia, research, and vendors is crucial to building the support network of resources that bring technology and design patterns across boundaries. Last year we focused on creating a community and demonstrating that there were existing approaches bridging these challenges. This session engages all participants in a discussion around infrastructure architecture design for a real problem to solve. Active collaboration builds lasting relationships, and by leveraging a working session we will solidify the community and create continued participation during the coming year.

Monitoring Large-Scale HPC Systems: Issues and Approaches

5:30pm-7pm

Room: 297

Jim Brandt (Sandia National Laboratories), Michael Showerman (University of Illinois at Urbana-Champaign), Michael Mason (Los Alamos National Laboratory)

This BoF addresses critical issues and approaches in large-scale HPC monitoring from the perspectives of system administrators, users, and vendors. In particular we target capabilities, gaps, and roadblocks in monitoring as we move to extreme scales including: a) desired information, b) vendor and tool-enabled interfaces to data, c) integration of capabilities

ties that provide and respond to data (e.g., integrated adaptive runtimes, application feedback), d) Monitoring impact analysis methods for large scale applications, and e) other hot topics (e.g., power, network congestion, reliability, high-density components). A panel of large-scale HPC stakeholders will interact with BoF attendees on topics of interest.

National Academies Study on Future Directions for NSF Advanced Computing Infrastructure to Support U.S. Science in 2017-2020

5:30pm-7pm

Room: 396

Jon Eisenberg (National Academy of Sciences), William Gropp (University of Illinois at Urbana-Champaign), Robert Harrison (Stony Brook University)

A National Academies study is examining priorities and associated tradeoffs for advanced computing in support of National Science Foundation-sponsored science and engineering research. A brief presentation of the study committee's interim report will be followed by comments and discussion on the issues raised in the report from the science and engineering communities that use, develop, and provide advanced computing capabilities.

OpenACC API User Experience, Vendor Reaction, Relevance, and Roadmap

5:30pm-7pm

Room: 275-76-77

Duncan Poole (NVIDIA Corporation), Barbara Chapman (University of Houston), James Bayer (Cray Inc.)

OpenACC API for accelerated systems has been earning praise for leadership in directives programming models and early implementations, and concerns from those looking for a single unified code for all system architectures. OpenACC SC13 BoF saw the OpenACC 2.0 announcement, received valuable input from the user community that helped us drive improvements to the specification and implementation. In this BoF, we will examine key user questions, encourage dialog between developers and implementers, discuss the current and future status of the specification, help developers contrast several existing accelerator programming solutions, highlighting the choices for given science domains, among other topics.

OpenSHMEM: Further Developing the SHMEM Standard for the HPC Community

5:30pm-7pm

Room: 294

Tony Curtis (University of Houston), Steve Poole, Oscar Hernandez (Oak Ridge National Laboratory)

The purpose of this BOF is to engage collaboration and input from the HPC community to further expand the OpenSHMEM specification with the help of users, hardware vendors, and tools developers. As a result of last year's BOF, we released the OpenSHMEM 1.1 specification. At this year's BOF, we will announce the upcoming OpenSHMEM 1.2 specification. We will also present any new developments in the different OpenSHMEM implementations, tools, extensions, and any novel new hardware support. We also plan to discuss the new OpenSHMEM roadmap for exascale. The BOF is an excellent opportunity to provide input into this ongoing process.

Programmable Storage Systems

5:30pm-7pm

Room: 398-99

Carlos Maltzahn (University of California, Santa Cruz), Brent Welch (Google), Pat McCormick (Los Alamos National Laboratory)

With the advent of open source storage systems a new usage pattern emerges: users are reusing subsystems and put them in contexts not foreseen by original designers. However, these attempts often prove to be more difficult than they should be. This BoF focuses on frameworks that allow the programmability of file and storage systems to enable unprecedented optimizations while preventing data corruption. It aims to serve as a venue for leaders in the file system and storage community to exchange ideas outside the tradition of classic file and storage systems research.

Super-R: Supercomputing and R for Data-Intensive Analysis

5:30pm-7pm

Room: 388-89-90

Weijia Xu (University of Texas at Austin), Hui Zhang (Indiana University), George Ostrouchov (Oak Ridge National Laboratory)

R has become popular for data analysis in many areas due to its high-level expressiveness and multitude of domain-specific packages. Challenges still remain in developing methods to effectively scale R to the power of supercomputers, and in deploying and enabling access to end users. This year's BOF will focus on usability and provisioning R as an interactive high performance analysis environment within HPC resources. This BOF will consist of a set of short presentations and discussions

with audience. The ultimate goal is to help build a community of users and experts interested in applying R to solve data intensive problems.

The Eclipse Parallel Tools Platform

5:30pm-7pm

Room: 291

Beth Tibbitts (University of Illinois at Urbana-Champaign), Greg Watson (IBM)

The Eclipse Parallel Tools Platform (PTP) (<http://eclipse.org/ptp>) is an open-source project providing a robust, extensible workbench for the development of parallel and scientific codes. PTP makes it easier to develop, build, run, optimize, and debug parallel codes on a variety of remote clusters using a single unified interface. PTP includes support for MPI, OpenMP, UPC, Fortran, and other libraries as well.

This BOF will consist of brief demos and discussions about PTP, and an overview of upcoming features. Information from contributors and vendors is welcome concerning integrations with PTP. How to get involved in the PTP project will be covered.

The HPCG Benchmark: Getting and Interpreting Performance from a New Metric for HPC Systems

5:30pm-7pm

Room: 273

Michael Heroux (Sandia National Laboratories), Piotr Luszczek (University of Tennessee, Knoxville)

The High Performance Conjugate Gradients (HPCG) Benchmark is emerging as a new community metric for ranking high performance computing systems. The first list of results was released at ISC'14, including optimized results for systems built upon Fujitsu, Intel, NVIDIA technologies.

In this BOF we first present the organization of HPCG and opportunities for optimizing performance and follow with presentations from vendors and leadership computing facilities who have participated in HPCG optimization efforts. We spend the remaining time in open discussion about the future of HPCG design and implementation strategies for further improvements.

The Message Passing Interface: MPI 3.1 and Plans for MPI 4.0

5:30pm-7pm

Room: 293

Martin Schulz (Lawrence Livermore National Laboratory)

The MPI forum, the standardization body for the Message Passing Interface (MPI), is in the process of preparing version 3.1 of the MPI standard, which will feature minor

improvements and corrections over MPI 3.0 released two years ago. Concurrently, the MPI forum is working on the addition of major new items for a future MPI 4.0, including mechanisms for fault tolerance, improved support for hybrid programming models, and streaming communication. We will use this BoF to introduce the changes made in MPI 3.1 and to continue an active discussion with HPC community on features and priorities for MPI 4.0.

Xeon Phi Users Group: Performance Tuning and Functional Debugging for Xeon Phi

5:30pm-7pm

Room: 286-87

Richard Gerber (National Energy Research Scientific Computing Center), Kent Milfeld (University of Texas at Austin), Chris Gottbrath (Rogue Wave Software)

This BOF will build community among those developing HPC applications for systems incorporating the Intel Xeon Phi many-core processor. Taking advantage of the processor's full capabilities requires tuning and optimizing using programming techniques and tools targeted at a combination of CPUs and the Xeon Phi Coprocessor. Threading, vectorization, memory contiguity and alignment, and data locality may be important for performance.

This BOF will combine a few brief presentations sharing insights and best practices with a moderated discussion among all those in attendance. It will close with an invitation to an ongoing discussion through the Intel Xeon Phi Users Group (IXPUG).

Thursday, November 20

Analyzing Parallel I/O

12:15pm-1:15pm

Room: 286-87

Julian Kunkel (German Climate Computing Center), Philip Carns (Argonne National Laboratory), Alvaro Aguilera (Technical University Dresden)

Parallel application I/O performance often does not meet user expectations. Moreover, slight access pattern modifications may lead to significant changes in performance due to complex interactions between hardware and software. These challenges call for sophisticated tools to enable scientists and facilities operators to capture, analyze, understand, and tune application I/O.

In this BoF we will introduce tools that can help to address this problem for different use cases. We focus in particular on the features of Darshan, SIOX, and Vampir. We also seek to gather feedback from the community regarding I/O challenges and discuss the requirements for monitoring future systems.

Charm++: Adaptive and Asynchronous Parallel Programming

12:15pm-1:15pm

Room: 291

Laxmikant Kale, Eric Bohm (University of Illinois at Urbana-Champaign)

A BoF for the Charm++ parallel programming system and the associated ecosystem (Adaptive MPI, mini-languages, tools etc.) is planned. The session will include discussion of the recent advances of the Charm++ system, and the major applications developed using it. Furthermore, there will be discussion about the future directions of Charm++. The goal is to engage a broad audience, drive adoption, and plan future developments.

Interest in Charm++ is increasing. It is one of the main programming systems deployed on parallel supercomputers, and a large fraction of processor cycles are devoted to Charm++ applications (such as NAMD).

Computational and Data Challenges in Genomic Sequencing

12:15pm-1:15pm

Room: 273

Patricia Kovatch (Icahn School of Medicine at Mount Sinai), Dan Stanzione (University of Texas at Austin), Shane Canon (Lawrence Berkeley National Laboratory)

As genomic sequencers proliferate, more HPC centers are assisting with genomic workflows, which are very different from traditional HPC workloads. Over time, these workflows have become vastly more complicated with the diversification of genomic sequencers, and their cumulative data footprint has risen significantly. To tackle the rising complexities, our goal is to build a community of scientists and infrastructure specialists to share common experiences. Visit <http://sc14bof.hpc.mssm.edu> to input your questions, vote for BOF discussion topics and answer questions from others. Our dialogue will be guided by real-time feedback gathered from this website. Engagement after SC will be explored and established.

Configuration Management Tools for Large Scale HPC Clusters

12:15pm-1:15pm

Room: 388-89-90

Sharan Kalwani, Amitoj Singh (Fermi National Accelerator Laboratory)

Open Source or proprietary configuration management tools? Tons of choices out there, this is geared towards the much harried sysadmin for any size HPC cluster and the nightmare of keeping everything in sync and up to date. What choices are out there? How to pick one over the other? Come and find out!

Defining Interfaces for Interoperable Simulation and Modeling Tools

12:15pm-1:15pm

Room: 295

Bruce Childers (University of Pittsburgh), Noel Wheeler (Laboratory for Physical Sciences / Advanced Computing Systems), Daniel Mosse (University of Pittsburgh)

This BOF aims to engage academia, government and industry in discussion of open frameworks and repositories for repeatable experimentation with interoperable and validated simulators for computer architecture. This BOF will solicit SC attendees to get involved and provide feedback in defining interfaces for interoperability and creating the tools that use the interfaces. It will help build a community behind interoperable, open simulation and modeling. Attendees will be informed and update about an ongoing NSF project: Open Curation for Computer Architecture Models (OCCAM), which is developing a repository of curated interoperable simulation tools and experimental results gathered with the tools.

Drilling Down: Understanding User-Level Activity on Today's Supercomputers with XALT

12:15pm-1:15pm

Room: 386-87

Mark Fahey (University of Tennessee, Knoxville), Robert McLay (University of Texas at Austin)

Let's talk real, no-kiddin' supercomputer analytics: drilling down to the level of individual batch submissions, users, and binaries. And we're not just targeting performance: we're after everything from which libraries and functions are in demand to preventing the problems that get in the way of successful science. This BoF will bring together those with experience and interest in technologies that can provide this type of job-level insight. To this end, the proposers will show off their new tool named XALT and have a short demo that the attendees can do on their own laptops.

HPC Centers for Global Health

12:15pm-1:15pm

Room: 398-99

Ian Brooks (University of Illinois at Urbana-Champaign)

Many HPC Centers around the world are under increasing pressure to show societal benefit and are looking to healthcare as a potential area in which to make a significant contribution. Following on from the Global Health session at the recent el4Africa conference hosted by the Tanzanian HPC Center this, session will feature an open discussion on HPC in healthcare and how we as a community can make a difference. We propose to use this BoF session to begin developing a network of HPC Centers that will work with WHO/PAHO to apply their collective resources for global health improvement.

Managing High-Performance Computing Centers: Communicating the Value of HPC

12:15pm-1:15pm

Room: 393-94-95

Nicholas Berente (University of Georgia), John King (University of Michigan), Daniel Reed (University of Iowa)

It can be difficult to communicate the value of HPC to those outside of the HPC community. The goal of this birds-of-a-feather session is to discuss different ways that HPC leaders can communicate the value of HPC to different stakeholders, and to try to uncover lessons learned and best practices.

The session will begin with opening statements by panelists with diverse experience, including industry, academic administration, government, and HPC leadership. This will be followed by general discussion.

OpenStack Technology and HPC Cloud Solutions

12:15pm-1:15pm

Room: 297

Thomas Goepel (Hewlett-Packard), Glen Luft (Hewlett-Packard)

This session will be an interactive discussion with panel leaders who will share their visions, work, and research with OpenStack technology in HPC solutions to date. They will also discuss with attendees the issues and opportunities that OpenStack technology offers to create open, HPC cloud environment that can be optimized to meet HPC workload needs and more easily integrated with other cloud infrastructures. If you are interested in exploring, participating, or sharing what you have learned to help in advancing OpenStack usage in HPC solutions we invite you to join us for an enabling exchange of ideas.

Propose and Define an Abstracted Interface for General Usage

12:15pm-1:15pm

Room: 396

Dave Resnick (Sandia National Laboratories), Mike Ignatowski (Advanced Micro Devices, Inc.)

It is getting increasingly obvious that current computer system interfaces need more capabilities at the same time that they need generalization. There will be network interfaces that have intelligence built in and memory components that have many more capabilities than current 'dumb' memory parts. An interface and protocol must be defined to enable this in such a way that it is very general so that many different kinds of current and future components can be interconnected.

Here is a paper that discusses the needs for and the benefits of a new interface and indicates some possible directions: https://sst-simulator.org/downloads/abstract_interface_final.pdf

SLURM User Group

12:15pm-1:15pm

Room: 283-84-85

Morris Jette (SchedMD LLC), David Wallace (Cray Inc.), Eric Monchalin (Bull)

SLURM is an open source job scheduler used on roughly half of Top500 systems and provides a rich set of features including topology aware optimized resource allocation, the ability to expand and shrink jobs on demand, the ability to power down idle nodes and restart them as needed, hierarchical bank accounts with fair-share job prioritization and many resource limits. The meeting will consist of three parts: The SLURM development team will present details about changes in the new version 14.10, describe the SLURM roadmap, and solicit user feedback. Everyone interested in SLURM use and/or development is encouraged to attend.

System Testing and Resiliency in High Performance Computing

12:15pm-1:15pm

Room: 391-92

Ashley Barker (Oak Ridge National Laboratory), Tim Robinson (Swiss National Supercomputing Center), Frank Indiviglio (National Oceanic and Atmospheric Administration)

The verification of hardware components and software updates is an ever more important aspect of system management as HPC systems become increasingly larger and more complex. Various custom frameworks for regression testing have been developed and deployed at different HPC centers to detect errors or performance regressions, allowing centers

to monitor issues ranging from run-time variability to bitwise reproducibility; however, to date there has been very little collaboration between centers towards these goals. This session will bring together those with experience and interest in regression testing theory and practice, with the aim of fostering collaboration and coordination across centers.

The Future of Fortran

12:15pm-1:15pm

Room: 293

Steven Lionel (Intel Corporation)

Join Fortran standards committee members and industry experts for a look at the current state of Fortran 2015, the next Fortran standard. After a brief presentation, an open discussion among panelists and participants will follow.

The Green500 List and its Continuing Evolution

12:15pm-1:15pm

Room: 275-76-77

Wu-chun Feng (Virginia Polytechnic Institute and State University), Erich Strohmaier (Lawrence Berkeley National Laboratory), Natalie Bates (Energy Efficient HPC Working Group)

The Green500, entering its eighth year, encourages sustainable supercomputing by raising awareness in the energy efficiency of such systems. This BoF will present (1) evolving metrics, methodologies, and workloads for energy-efficient HPC, (2) trends across the Green500, including the trajectory towards exascale, and (3) highlights from the latest Green500 List. In addition, the BoF will discuss a collaborative effort between Green500, Top500, and EE HPC WG to improve power measurement while running a workload, such as HPL, and solicit feedback from the HPC community. The BoF will close with an awards presentation, recognizing the most energy-efficient supercomputers in the world.

The Open Community Runtime (OCR) Framework for Extreme Scale Systems

12:15pm-1:15pm

Room: 294

Vivek Sarkar (Rice University), Barbara Chapman (University of Houston), William Gropp (University of Illinois at Urbana-Champaign)

Extreme-scale and exascale systems impose new requirements on software to target platforms with hundreds of homogeneous and heterogeneous cores, as well as energy, data movement and resiliency constraints within and across nodes. The Open Community Runtime (OCR) was created to engage the broader community of application, compiler, runtime, and hardware experts to co-design a critical component of the future exascale hardware/software stack. We will build on

the well-attended OCR BOFs in SC12 and SC13, by soliciting feedback on new extensions to OCR for intra- and inter-node parallelism and on the new Modelado foundation open source hosting site for OCR.

The Square Kilometer Array: Next Generation Sensor Networks as a Driver for HPC

12:15pm-1:15pm

Room: 383-84-85

Chris Broekema (Netherlands Institute for Radio Astronomy),

Rob Simmonds (University of Calgary)

The Square Kilometer Array (SKA) radio telescope is a next generation sensor network to be built in western Australia and southern Africa from 2017. Designing the exascale computer systems required to process data from such instruments presents challenges due to huge amounts of data, high data rates, relatively low computational intensity and a limited power budget. Such computer systems can be considered data throughput machines applicable to a range of large scale big data processing tasks. The SKA Science Data Processor consortium invites the community to explore opportunities for collaboration, and offers an update of its current design status.



Doctoral Showcase

Doctoral Showcase

This program provides an important opportunity for students near the end of their Ph.D. to present a summary of their dissertation research in the form of short talks and posters. Unlike technical paper and poster presentations, Doctoral Showcase highlights the entire contents of each dissertation, including previously published results, to allow for a broad perspective of the work.

The SC14 Doctoral Showcase received 44 submissions that underwent rigorous peer review. Each submission received a minimum of three reviews from a committee with expertise across a wide range of technical areas in high performance computing, networking, storage and analysis. As a result of this competitive process, 16 students were selected to present their research in the Showcase.

We invite you to see the latest work by some of our brightest rising stars!

Doctoral Showcase

Thursday, November 20

Doctoral Showcase

Chair: Karen L. Karavanic (Portland State University)

Dissertation Research I

10:30am-12pm

Room: 386-87

Fast Storage for File System Metadata

Kai Ren (Carnegie Mellon University)

Conventional distributed have used centralized single-node metadata services. However, the single-node metadata server design inherently limits the scalability of the file system in terms of the number of stored objects and concurrent accesses to the file system expected by massive parallel applications. Inefficient on-disk metadata representation also limits the metadata performance. To tackle these challenges, I implemented a middleware called IndexFS that can be layered on top of existing file systems and improve their metadata performance. IndexFS uses a tabular-based architecture that incrementally partitions the namespace on a per- directory basis, preserving server and disk locality for small directories. An optimized log-structured layout is used to store metadata and small files efficiently. Several caching techniques are also used to mitigate hot spots. By combining these techniques, IndexFS can improve the metadata performance of existing distributed file systems by as much as an order of magnitude for various metadata workloads.

Designing Scalable and Efficient I/O Middleware for Fault-Resilient High-Performance Computing Clusters

Raghunath Raja Chandrasekar (Ohio State University)

This dissertation proposes a cross-layer framework that leverages the hierarchy in storage media (memory, ramdisk, flash/ NVM, disk, parallel FS, and so on), to design scalable and low-overhead fault-tolerance mechanisms that are inherently I/O bound. The key components of the framework include - CRUISE, a highly-scalable in-memory checkpointing system that leverages both volatile and Non-Volatile Memory technologies; Stage-FS, a light-weight data-staging system that leverages burst-buffers and SSDs to asynchronously move application snapshots to a remote file system; Stage-QoS, a file system agnostic Quality-of-Service mechanism for data-staging systems that minimizes network contention; MIC-Check, a distributed checkpoint-restart system for coprocessor-based supercomputing systems; and FTB-IPMI, an out-of-band fault-prediction mechanism that pro-actively monitors for failures.

Storage Support for Data-Intensive Applications on Extreme-Scale HPC Systems

Dongfang Zhao (Illinois Institute of Technology)

Many believe that current HPC design would not meet the I/O requirement of the emerging exascale computing due to the segregation of compute and storage resources. Indeed, our simulation predicts, quantitatively, that system availability would go towards zero at exascale. This work proposes a storage architecture with node-local disks for HPC systems. Although co-locating compute and storage is not a new idea, it has not been widely adopted in HPC systems. We build a node-local filesystem, FusionFS, with two major principles: maximal metadata concurrency and optimal file write, both of which are crucial to HPC applications. We also discuss FusionFS' integral features such as hybrid and cooperative caching, efficient accesses to compressed files, space-efficient data redundancy, distributed provenance tracking, and integration with data management systems. We have evaluated FusionFS on petascale supercomputers with 64K-cores and, compared its performance with major storage systems such as GPFS, PVFS, HDFS, and S3.

High Performance Computing over Geo-Spatial Data: A Summary of Results

Satish Puri (Georgia State University)

Polygon overlay and polygon clipping are one of the complex operations in Geographic Information Systems (GIS), Computer Graphics, and VLSI CAD. We present the first output-sensitive parallel algorithm, which can perform geometric intersection, union, and difference operations in $O(\log n)$ time using $O(n + k)$ processors, where n is the number of polygonal vertices and k is the number of edge intersections. This is cost-optimal when compared to a sequential plane-sweep based algorithm and superior to the $O(n^2 \log n)$ time parallel algorithm by Karinthi, Srinivas, and Almasi. We have also developed two systems: (1) MPI-GIS and (2) Hadoop Topology Suite for distributed polygon overlay using a cluster of nodes. Our MPI and GPU-based system achieves 44X speedup while processing about 600K polygons in two real-world GIS shapefiles (1) USA Detailed Water Bodies and (2) USA Block Group Boundaries) within 20 seconds on a 32-node (8 cores each) IBM iDataPlex cluster interconnected by InfiniBand technology.

Enabling Scalable Data Analysis of Computational Structural Biology Datasets on Distributed Memory Systems Supported by the MapReduce Paradigm

Boyuan Zhang (University of Delaware)

In this research project, we focus on scalable and accurate classification and clustering analyses of large computational structure biology datasets on large distributed memory systems. We propose a transformative data analysis method that comprises of two general steps. The first step extracts concise properties or features of each data record in parallel and represents them as metadata. The second step performs the analysis (i.e., classification or clustering) on the extracted properties. Our method naturally fits in the MapReduce paradigm; we adapt it for different MapReduce frameworks (i.e., Hadoop, MapReduce-MPI, and DataMPI). We use the frameworks for three scientific datasets of RNA secondary structures, ligand conformations, and folding proteins. The evaluation results show that our method can perform scalable classification and clustering analyses on large-scale datasets that are generated and stored in a distributed manner. Moreover, our method achieves better accuracy comparing to the traditional approaches.

Accelerating MPI Collective Communications through Hierarchical Algorithms with Flexible Inter-Node Communication and Imbalance Awareness

Benjamin Parsons (Purdue University)

This work investigates collective communication algorithms on a shared memory system, and develops the universal hierarchical algorithm. This algorithm can pair arbitrary hierarchy unaware inter-node communication algorithms with shared memory intra-node communication. In addition to flexible inter-node communication, this algorithm works with all collectives, including those incompatible with past works, like all-to-allv. The universal algorithm shows impressive performance results, improving upon the MPICH algorithms as well as the Cray MPT algorithms. Speedups average 15-30x for most collectives with improved scalability up to 64k cores.

The second part of this work creates new hierarchical collective algorithms designed to tolerate process imbalance. The process imbalance of benchmarks is thoroughly evaluated and is used to design collective algorithms that minimize the synchronization delay observed by early arriving processes. Preliminary results for a reduction show speed-ups reaching 47x over a binomial tree algorithm in the presence of high, but not unreasonable, imbalance.

Dissertation Research II

1:30pm-3pm

Room: 386-87

Failure Avoidance Techniques for HPC Systems Based on Failure Prediction

Ana Gainaru (University of Illinois at Urbana-Champaign)

As the size of HPC systems continues to increase, so does the probability of single component failures within a time frame. With MTBFs of less than one day to a few hours for future systems, current fault tolerance strategies face serious limitations. My research focuses on offering ways of reducing the overhead induced by these strategies, by combining them with failure avoidance methods. My key observation that errors are often predicted by changes in the frequency or regularity of various events has encouraged the use of signal analysis concepts for failure analysis. The results on various systems show that signal processing techniques can create clear markers for changes in events behavior. Moreover, machine learning techniques become much more efficient when applied to the derived markers, rather than the original signal. Hybrid fault tolerance strategies based on this novel predictor applied on various systems have shown overhead reductions of over 20%.

Introspective Resilience for Exascale High Performance Computing Systems

Saurabh Hukerikar (University of Southern California)

Future exascale HPC systems will be constructed using hundreds of millions of components organized in complex hierarchies to satiate the demand for faster and more accurate scientific computations. However, the sheer scale and inherent unreliability of the VLSI chips implies that faults and errors will affect HPC applications with increasing frequency, making it increasingly difficult to accomplish useful computation. In this dissertation work, we propose an introspective approach to managing the resilience of HPC applications. Through a set of modest extensions to current programming models we capture the programmer's insight into the application's fault tolerance features. An introspective runtime framework reasons about the rate and sources of faults and errors in the system and attempts to understand their impact on application correctness. This enables collaborative cross-layer efforts for error detection and recovery. Our preliminary results demonstrate much promise to meet the reliability demands and expectations of future exascale platforms.

Fault Tolerant Iterative Solvers through Selective Reliability and Skeptical Programming

James Elliott (North Carolina State University)

Problem: Current and future systems may experience transient faults that silently corrupt data being operated on. These faults are troubling, because there is no indication a fault occurred and the cost to detect an error may involve performing operations multiple times and voting to determine if any values are tainted.

Approach: I focus on numerical methods and have proposed an approach called Skeptical Programming. My research couples the numerical method and the properties of data being operated on to derive cheap invariants that filter out large “damaging” errors, while allowing small “bounded” errors to slip through. The errors that slip through are easily handled by convergence theory, resulting in a low-overhead algorithm-based fault tolerance approach that exploits both numerical analysis and system-level fault tolerance. This technique scales well since we do not require additional communications, and is applicable to many methods (we present findings for the CG and GMRES solvers).

Semantically Ordered, Parallel Execution of Multiprocessor Programs

Gagan Gupta (University of Wisconsin-Madison)

Conventional wisdom says that we must abandon the logical order between computations to obtain high performance from multiprocessor programs. Unfortunately, the resulting nondeterminism complicates parallel programming, resource management, recovery from exceptions, and parallel system design in general. Conversely, uniprocessor systems have long benefited from ordered programs. Order leads to deterministic execution, which does not suffer from the shortcomings of nondeterminism. This work shows that an ordered approach can also be applied to multiprocessor programs, simplifying parallel system design, but without compromising performance.

Reliability and Energy Data Analysis and Modeling for Extreme Scale Systems

Li Yu (Illinois Institute of Technology)

Reliability and energy are two major concerns in the development of today’s supercomputers. To build a powerful machine while at the same time satisfying reliability requirements and energy constraints, HPC scientists continue to seek a better understanding of system and component behaviors. Toward this end, modern systems are deployed with various monitoring and logging tools to track reliability and energy data during system operations. Since these data contain important infor-

mation about system reliability and energy, they are valuable resources for understanding system behaviors. However, as system scale and complexity continue to grow, the process of collecting system data to extracting meaningful knowledge out of overwhelming reliability and energy data faces a number of key challenges. To address these challenges, my work consists of three parts, including data preprocessing, data analysis and advanced modeling.

Adaptive Power Efficiency: Runtime System Approach with Hardware Support

Ehsan Totoni (University of Illinois at Urbana-Champaign)

Power efficiency is an important challenge for the HPC community, demanding major innovations. We propose novel techniques using a cross-layer approach to improve power efficiency. We use extensive application-centric analysis of different architectures to design automatic adaptive runtime system (RTS) techniques that save significant power. These techniques exploit common application patterns and only need minor hardware support. The application’s pattern is recognized using Formal Language Theory to predict its future and adapt the hardware appropriately.

I will discuss why some system components such as caches and network links consume extensive power disproportionately for common HPC applications, and how a large fraction of power consumed in caches and networks can be saved using our approach automatically. In these cases, the hardware support the RTS needs is the ability to turn off ways of set-associative caches and network links.

Dissertation Research III

3:30pm-5pm

Room: 386-87

Parallel Algorithms for Two-Stage Stochastic Optimizations

Akhil Langer (University of Illinois at Urbana-Champaign)

Many real-world planning problems require searching for an optimal solution in the face of uncertain input. In a two-stage stochastic optimization problem, search for an optimum in one stage is informed by the evaluation of multiple possible scenarios in the other stage. Stochastic optimization has applications in production, financial modeling, transportation (road as well as air), supply chain, scheduling, environmental and pollution control, telecommunications and electricity. In this work, we explore the parallelization of a two-stage stochastic integer program solved using branch-and-bound. We present a range of factors that influence the parallel design for such problems. Unlike typical, iterative scientific applications, we encounter several interesting characteristics that make it challenging to

realize a scalable design. We evaluate the scalability of our designs and compare its performance with the state-of-the-art integer linear program solvers such as Gurobi. Our attempts result in strong scaling to hundreds of cores for these datasets.

Scalable, Adaptive Methods for Forward and Inverse Modeling of Continental-Scale Ice Sheet Flow

Tobin Isaac (University of Texas at Austin)

Projecting sea-level rise is made difficult by the complexity of accurately modeling ice sheet dynamics for the polar ice sheets and the uncertainty in key, unobservable parameter fields; my research addresses the inference of the basal friction field beneath the Antarctic ice sheet. I develop scalable algorithms and numerical methods that make tractable the calculation of a friction field with quantified uncertainties. These contributions fall in the categories of adaptive mesh refinement (AMR), efficient solvers for nonlinear PDEs, and Bayesian statistical inversion, all with an emphasis on scalability and high performance computing. I have developed algorithms for octree-based AMR that have scaled well to 458K processes on ~30K BG/Q nodes. I have developed a solver for high-order discretizations of the nonlinear Stokes equations of ice sheet dynamics that scales to ~600M dofs. I am developing Hessian-approximation techniques for Bayesian inference for problems whose parameter-to-observable map requires solving systems of PDEs.

Load Balancing Scientific Applications

Olga T. Pearce (Texas A&M University)

Optimizing high-performance physical simulations to run on ever-growing supercomputing systems is challenging. The dynamic behavior of large modern parallel simulation codes can lead to imbalances in computational load among processors. Load imbalance is particularly expensive at scale, because hundreds of thousands of idle processors may wait on a single overloaded processor. Future machines will support even more parallelism and efficiently redistributing and balancing load will be critical for good performance.

In this thesis, I address how to evaluate load imbalance and make its correction affordable. I developed a model for comparison of load balance algorithms in the context of a specific application imbalance scenario, enabling the selection of a balancing algorithm that will minimize overall runtime. I provide an accurate and fast method to balance the load in simulations with highly non-uniform density. I devised a framework for decoupling the load balancer from the application, enabling asynchronous load balancing.

Hierarchical Scheduling Frameworks for Heterogeneous Clusters with GPUs

Kittisak Sajjapongse (University of Missouri)

Graphic Processing Units have increasingly been adopted for a wide-range of scientific applications, and have become part of HPC clusters. Distributed GPU applications typically offload computation to GPUs using CUDA and OpenCL and distribute tasks through MPI and SHMEM. Despite the availability of these frameworks, coding such applications is still non-trivial. In addition, the use of batch schedulers to handle these applications on shared clusters often leads to performance and underutilization issues.

In our research, we propose the design of hierarchical scheduling frameworks consisting of node- and cluster-level runtime to support concurrent distributed GPU applications. Our proposed node-level runtime enables GPU virtualization, sharing and flexible scheduling mechanisms. The cluster-level scheduler allows administrator to define scheduling policies and configure node-level sharing. Results show that our framework outperforms existing batch schedulers while improving GPU utilization. Current work focuses on increasing the programmability of hybrid nodes while enabling effective sharing and load-balancing mechanisms.

Doctoral Showcase Research Poster Session

4:30pm-5pm

Room 386-87

Our Doctoral Showcase presenters, a select group of students each within a year of completing their HPC-related Ph.D., will be available in a poster session for discussion and questions about their work.



Exhibitor Forum

Exhibitor Forum

The Exhibitor Forum offers an opportunity to learn about the latest advances in the supercomputing marketplace. Manufacturers, vendors and HPC service providers will present new products, services and roadmaps. Industry leaders will share their insights into customer needs, market drivers, product strategies and technology trends. Case studies will illustrate how their current and future solutions can be used effectively.

Please join us for a few Exhibitor Forum sessions this year—we have a great lineup!

Exhibitor Forum

Tuesday, November 18

Hardware and Architecture

Chair: John Cazes (University of Texas at Austin)

10:30am-12pm

Room: 292

Interconnect Your Future

Gilad Shainer, Scot Schultz (Mellanox Technologies)

The exponential growth in data and the ever growing demand for higher performance to serve the requirements of the leading scientific applications, drive the need for Petascale systems and beyond and the ability to connect tens-of-thousands of compute and co-processor nodes in a very fast and efficient way. The interconnect has become the enabler of data and the enabler of efficient simulations. Beyond throughput and latency, the data center interconnect needs be able to offload the processing units from the communications work in order to deliver the desired efficiency and scalability. Mellanox has already demonstrated 100Gb/s cable solutions in March 2014 and announced the world first 100Gb/ switch at the ISC'14 conference, June 2014. The presentation will cover the latest technology and solutions from Mellanox that connect the world fastest supercomputers, and a roadmap for the next generation InfiniBand speed.

E4-ARKA: ARM64+GPU+IB is Here

Piero Altœ (E4 Computer Engineering)

E4 Computer Engineering introduces ARKA, the first ever server solution based on ARM 64 bit SoC dedicated to HPC. The compute node is boosted by discrete GPU NVIDIA cards K20, and have both 10Gb ethernet and FDR infiniband networks implemented by default. The CPU is an APM X-gene with 4 memory channel and can support up to 128 GB of DDR3 ECC memory. In the ARKA machine 8 PCIe gen3 lines are connected to the GPU and other 8 lines to the FDR infiniband card. A dual port 10GbE SFP+ is embedded on the SoC. The remote management is guaranteed by an BMC module IPMI 2.1 compatible. The 1U short depth form factor is the best compromise between density and costs, and fits in all datacenter and research labs. E4-ARKA is the perfect server to experience the ARM revolution in a professional solution.

KALRAY TURBOCARD2 Scalable Accelerator: An Efficient Supercomputer Cluster on a Board

Benoît Ganne (Kalray Corporation)

Kalray, a fabless semiconductor company, has developed an innovative on-chip supercomputer technology that provides more processing performance without increasing power consumption. Kalray's processor, the 256 cores MPPA-256, delivers up to a quarter of a TeraFLOPS and offers one of the top GFLOPS / Watt ratio available on the market today.

At SC14 Kalray will announce the availability of KALRAY TURBOCARD2, which brings 1024 cores on a PCI-e card with unique power efficiency. It also includes up to 32 GB of DDR3 and High throughput Low latency Network on Chip Extension (NoCX) fabric to connect card to card.

The presentation will focus on the architectural aspects of the MPPA processor and the new TURBOCARD2, its software development environment and its ability to scale up and provide high density computing solutions for a wide spectrum of application domains such as video compression, cryptography, finance and oil & gas.

Moving, Managing, and Storing Data

Chair: Dane Skow (Self)

10:30am-12pm

Room: 291

Why Archive? Drive Down Costs and Improve Efficiencies

Christine Rogers (Oracle Corporation)

Data growth, regulations, content digitization, lean budgets, and the need for storage efficiencies are all complicating the operational and financial pictures of enterprises, which must store and protect information for years. While data retention challenges are growing, organizations are also looking to leverage their data to gain insight into and grow their businesses.

Learn how introducing an archive strategy for long-term data retention can help your organization drive down storage costs and improve operational efficiencies. Your organization can, then, focus on unlocking the value of your data for new business opportunities.

Eliminating Data Junkyards with Metadata

Robert Murphy (General Atomics)

Data stored by many HPC organizations is increasing exponentially. And while today's high capacity scale-out file and object storage systems can easily accommodate the growing volume

of data, their limited, bare-bones system metadata can leave high-value data lost over time – stranding it and losing its value forever. Without application-specific metadata, HPC organizations will amass a growingly unmanageable “Data Junkyard” of inaccessible, ultimately useless data, at very high cost. This presentation will explore how the Big Data challenges faced by Big Science projects like CERN’s Large Hadron Collider, and the metadata strategies those projects pioneered, could be exploited to support data intensive HPC workflows by using the metadata management facilities of the Nirvana Metadata Centric Intelligent Storage system. Nirvana was developed by the General Atomics Magnetic Fusion Energy and Advanced Concepts group in San Diego, California from a joint effort with the San Diego Supercomputing Center’s Storage Resource Broker (SRB).

Scale Performance for File-Based Applications in an Instant with Cloud Bursting

Jeff Tabor (Avere Systems)

Organizations running high-performance, file-based applications in fields such as life sciences, oil & gas, and defense are under constant pressure to scale compute resources for ever-increasing performance demands. Cloud computing services like Amazon EC2 offer near infinite compute resources for such applications but introduce new challenges. How do I get my data to the compute cloud for processing? How do I avoid the latency of the WAN? How do I scale file system performance in the compute cloud?

Avere Systems addresses these challenges with a Hybrid Cloud NAS solution that is the first scale-out NAS solution purpose-built for the cloud. With Avere, you can achieve the following in the compute cloud: 1) Scalable NAS performance through clustering, 2) High availability with active/active failover, 3) Access data stored on premises or in the cloud, 4) Caching hides WAN latency, 5) Full NFS and SMB/CIFS support, 6) Simple software-only installation

Hardware and Architecture

Chair: Carlos Rosales (University of Texas at Austin)

1:30pm-3pm

Room: 292

The George Washington University Addresses Extreme Heat in the Data Center

Tim Wickberg (George Washington University),

Patrick Giangrosso (Coolcentric)

Adding an HPC cluster to the existing data center at the George Washington University brought a serious concern to the forefront. Additional heat loads from the 24 kW per rack cluster proved too much for the existing hot aisle chimney strategy.

The 50°F cold aisle air was barely adequate and high back pressure at the back of the rack raised serious concerns that the chimneys would not be able to remove the extreme heat while maintaining the health of the HPC cluster.

To address cooling demands and future growth requirements, GW implemented Coolcentric’s High Density Rear Door Heat Exchangers (RDHx-HD); a passive heat exchanger closely coupled to the rear of the rack capable of removing 30kW of heat before it enters the data center.

Giangrosso and Wickberg will discuss GW’s decision, installation, results and implications for existing data center facilities that are at or near cooling and power capacities.

Powering Discovery and Insight with Accelerated Computing
Han Vanholder (NVIDIA Corporation)

GPU-accelerated computing has become an indispensable driver for the advancement of science, technology and high performance analytics. As everything we do becomes more data intensive, accelerated computing is playing an increasingly critical role in solving the most complex computational problems, including uncovering patterns in the largest data sets, faster and more efficiently.

NVIDIA continues to revolutionize HPC with its Tesla accelerated computing platform, deeply integrating GPU accelerators, software, tools and services to deliver unrivaled performance. Attend this talk to learn how GPU-accelerated computing powers discovery and insight across domains.

EXTOLL - The Interconnect Component for Dense Supercomputing

Ulrich Bruening (EXTOLL GmbH)

EXTOLL is a high-performance interconnection network targeted for dense supercomputing. EXTOLL’s unique feature of a PCIe root-complex allows for autonomous boot and operation of accelerators without a direct connected host. Combining an accelerator like the Xeon Phi and the EXTOLL NIC into a node builds a very dense compute node which is called Booster. This host-less Booster concept was developed in the EU-Project DEEP. Aggregating such nodes in a 3-D torus topology allows building very dense compute cabinets. Now power density becomes an issue, which is solved at the cabinet level by a 2-phase immersion cooling. More than 10kWatt of power dissipation can be removed from the 32 Booster nodes providing high performance of 32 TFLOPS at a yet unachieved density.

Moving, Manging, and Storing Data

Chair: Dane Skow (Self)

1:30pm-3pm

Room: 291

Performance Critical Data Center Analytics with 20 Gbps RapidIO Interconnect

Rick O'Connor (RapidIO)

As high performance analytics oriented applications become more important in the data center, there is momentum for solving these problems with low latency interconnect, using heterogeneous compute elements including: PowerPC, x86, ARM & DSP/FPGA. Systems using 20 Gbps RapidIO interconnect with 100ns latencies, no TCP termination overhead, scaling to 1000's of nodes in any topology are being developed for data center trials and are also being deployed in the supercomputing market delivering performance of 6.4 Gflops per watt. This presentation describes the RapidIO Data Center Compute and Networking CPU agnostic platform and related systems.

Using Erasure Codes to Revolutionize Data Reliability at Scale

Geoffrey Noer (Panasas)

As the need for scale-out data storage in HPC continues to grow at an exponential rate, traditional data protection and availability models are being stretched past their breaking points. This talk will contrast the principal difficulties associated with hardware RAID and legacy availability models with what is made possible by instead using erasure codes to deliver a revolutionary per-file, triple parity protection storage architecture.

Most importantly, this new approach allows data reliability to increase with the scale of the storage deployed, as opposed to decreasing as is normally expected. It also significantly improves data availability by changing from an all-up/all-down model to still maintaining partial availability in the rare case of storage systems experiencing multiple, simultaneous, drive failures.

The Rebirth of Tape: Tape's Changing Role and Ever-Increasing Advantages Over Disk are Leading to its Resurgence

Kevin Horn (Oracle Corporation)

Not only is the role of tape evolving to meet the colossal storage needs of cloud service providers and other emerging applications, tape storage density and performance is expected to continue to increase for many years to come. As we will discuss in this session, advances in disk storage density and performance have slowed considerably - a big reason why

some analysts predict disk technology will become increasingly rare as flash and tape become the preferred storage mediums for the majority of active and inactive data, respectively. In this session, find out what the industry is saying about the bright future of tape and learn the key benefits of incorporating tape technology into your backup and archive storage environments.

Hardware and Architecture

Chair: Richard Evans (Texas Advanced Computing Center)

3:30pm-5pm

Room: 292

High Performance SDN Data Planes

Bruce Gregory (Corsa Technology)

Pure SDN Data Planes are a new element of hardware that create a true separation between control planes and data planes in SDN networks to allow real global orchestration of your network. To date, it has been very difficult to implement a high performance at scale SDN network because the available data planes have had serious limitations. During this session, Corsa will review all the features an SDN data plane must have in order to meet the demands of a large, high speed SDN network. We will dive into the Corsa SDN data plane architecture to fully explain how our multiple tables and deep packet buffers work together to form a data pipeline that can be used for QoS in big data networks, as an SDN BGP gateway, a Label Edge Router, an MPLS router and various other functions. Come learn about pure SDN data planes and their important role.

Fault-Tolerant InfiniBand Fabric Management

Sven-Arne Reinemo (Fabriscale), Tor Skeie (Fabriscale)

Fabriscale is a start-up company that specializes in fabric management software with an emphasis on smart algorithms that simplify network configuration, management and routing. At SC14 we will announce our first product, the Fabriscale Fabric Manager. In the first release of the Fabriscale Fabric Manager for InfiniBand our center of attention is topology agnostic routing and fast fault-tolerance. When the network fails the consequences can be dramatic, but with Fabriscale you are protected against the negative consequences of any switch or link failure independent of your network topology. Any incidents are handled automatically, transparently, and quickly by the Fabriscale routing engine.

This presentation will present an overview of Fabriscale's unique routing and fault-tolerance capabilities, its use cases and benefits.

Practical considerations for HPC systems architecture

Jean-Pierre Panziera (Bull)

Designing HPC systems for optimal application performance requires a tight integration of the compute elements, the interconnect and the storage system while minimizing cost and energy consumption. This presentation describes Bull's approach to HPC systems architecture and illustrate the concepts on some recent installations. The compute nodes are based on a combination of CPUs or HPC accelerators depending on the application profile. The interconnect based on InfiniBand is configured with a compact fat-tree topology. The storage nodes which aggregate a huge capacity (PBs) are carefully placed into the network fabric in order to provide fast data access with full bandwidth (100s GB/s) and high availability. For lower cost, the platform integrates compute nodes and parts of the interconnect. The direct liquid cooling technology minimizes the energy consumption required for cooling and thus reduces operating costs. Finally, the key ideas guiding the current developments for Bull next generation system are explained.

Moving, Managing, and Storing Data

Chair: Paul Domagala (Argonne National Laboratory)

3:30pm-5pm

Room: 291

Building High IOPS Flash Array with Innodisk FlexiRemap Technology

Charles Tsai (Innodisk Corporation)

With rapid advancement of computing technologies, disk access has been identified as the next performance bottleneck. In recent years, storage appliances based on flash memory have been deemed as practical solutions to this bottleneck. However, high-end flash appliances are mostly built with proprietary hardware designs, aiming at particular scenarios in larger-scale data centers, and hence are barely affordable by enterprise and industry customers that are also deploying private clouds. Innodisk FlexiRemap technology, on the other hand, deals with the challenges of performance, data endurance, and affordability through innovations in software and firmware, creating a new category of flash-collaborative storage appliances that deliver sustained high IOPS, even for random writes. In this session, the speaker will elaborate on the challenges and address them with Innodisk FlexiArray storage appliance, a high-IOPS flash array built upon Innodisk FlexiRemap technology, providing a cost-effective alternative to customers with demands for high-speed data access.

BeeGFS – Fraunhofer's Fast and Easy Parallel Filesystem for Fast and Easy Computing

Jan Heichler (ThinkParQ GmbH)

The talk will present the basic concepts and advantages of Fraunhofer's parallel filesystem BeeGFS (formerly known as FhGFS). It provides extreme performance and scalability and is very easy to use. It enables everyone to do scale out storage to speed up computations. Additionally, performance measurements of the latest version will be included to show the excellent scaling of storage- and metadata performance.

Transform Large-Scale Science Collaboration

Rodney Wilson (Ciena Corporation)

Science data collection and storage is growing at an exponential pace. This enormous amount of data is most useful when easily shared among members of the global research community, so the network that carries this information must be designed with very large data transfers in mind. Standard campus networks are not optimized to support the movement of enormous big science-related data files between instruments, facilities, analysis systems, and scientists. To fully reach the potential discoveries related to big science, the network that interconnects big data must be optimized accordingly and this is why the emerging Science DMZ network architecture is gaining in popularity. This presentation will cover the Science DMZ architecture and how Ciena's recently announced 8700 Packetwave Platform is a natural fit to allow big science to reach its full potential in terms of collaboration and discovery.

Wednesday, November 19

HPC Futures and Exascale

Chair: Wesley Bland (Argonne National Laboratory)

10:30am-12pm

Room: 292

OpenMP API Version 4.0 | Announcing OpenMPCon

Michael Wong (OpenMP)

The OpenMP® Application Program Interface (API) supports shared-memory parallel programming in C, C++, and Fortran on a wide range of computer and operating systems. The OpenMP API provides a portable, scalable programming model that gives shared-memory parallel programmers a simple, flexible interface for developing parallel applications for a wide variety of platforms.

The OpenMP API version 4.0 includes accelerator support for heterogeneous environments; Enhanced tasking model with groupings, dependencies, and scheduling constraints; Thread affinity to support binding and to improve performance on non-uniform memory architectures; and more.

In 2015 we will launch the new annual OpenMPCon: A face-to-face gathering organized by the OpenMP community for the community. Enjoy keynotes, inspirational talks, and a friendly atmosphere that helps attendees meet interesting people, learn from each other, and have a stimulating experience. Multiple diverse technical tracks are being formulated that will appeal to anyone, from the OpenMP novice to the seasoned expert.

Quantum Computing: A New Resource for the Most Complex Problems

Colin Williams (D-Wave Systems)

Despite the incredible power of today's supercomputers, there are many computationally-intensive applications that can't be addressed effectively by conventional systems. Quantum computers have left the laboratory and offer the potential to help solve some of the most complex technical, commercial, scientific, and national defense problems that organizations face. In the near term quantum computing will complement HPC by providing unique computing resources for complex optimization problems, Monte Carlo / sampling applications and provide new capabilities to advance machine learning.

In this presentation, Dr. Williams will explain the basics of quantum computing, the current state and future potential of this revolutionary technology, and the kinds of applications that it will address.

NEC SX-ACE Vector Supercomputer and Future Direction

Shintaro Momose (NEC Corporation)

NEC gives an overview of the latest model vector supercomputer SX-ACE with its future direction in high-performance computing. NEC has launched SX-ACE by aiming at much higher sustained performance particularly for memory-intensive applications. It is designed based on the big core strategy targeting higher sustained performance. It provides both the world top-level single core performance of 64 GFlops and the world highest memory bandwidth per core of 64 GBytes/s. Four cores, memory controllers, and network controllers are integrated into a single CPU LSI, enabling the CPU performance of 256 GFlops and the memory bandwidth of 256 GBytes/s with a 64 GB memory capacity. In the target memory-intensive application areas, SX-ACE demonstrates both higher sustained performance and power efficiency, surpassing other supercomputers. NEC will also talk about the vision and concept of its future supercomputer products.

Software for HPC

Chair: Dane Skow (Self)

10:30am-12pm

Room: 291

Virtualization for Flexible HPC Environments

Shai Fultheim (ScaleMP)

High end virtualization for HPC environments is an extremely valuable technology which helps to reduce costs and create a more flexible HPC environment. By aggregating industry-standard systems into powerful, large scale SMPs, organizations can quickly react to various workloads. Applications can get full access to terabytes of memory and hundreds of compute cores in a simple and easy to use system. The latest technology is available to users and in addition they can get access to massive quantities of compute cores and memory. In addition to fully supported versions, a free version is available so that demanding users can experience how their applications can scale, without limiting the apps to small amounts of memory.

How to Scale In-Memory Real-Time Analytics

Atle Vesterkjær, Laurence Liew (Numascale AS)

Atle Vesterkjær will discuss how big Numascale shared memory systems can be used to get excellent scaling of applications that have been hard to scale on traditional clusters. He presents methods to take advantage of the extreme number of cores that the Numascale systems enable. The combination of large memory and a high number cores make the Numascale Shared Memory solution more balanced for Big Data applications than any other large shared memory system. Applications in the field In-Memory Real-Time Analytics are especially addressed with emphasis on software that Numascale has implemented and offers on its NumaQ systems.

Laurence Liew will share his perspective on the current hype around big data and why people experienced with HPC are best positioned to ride this wave. He will discuss where the big data market is moving, and in particular the next-generation of in-memory analytics infrastructure and software tools.

HPC Compilers for Heterogeneous Supercomputing

Michael Wolfe (Portland Group)

A fundamental shift is under way in the HPC market. A diverse array of CPUs and Accelerators are feeding a de facto node architecture with high-speed latency-optimized SIMD-capable CPU cores coupled to massively parallel bandwidth-optimized Accelerators with exposed memory hierarchies. This presents challenges and opportunities for HPC compilers. To support and drive innovation, compilers must now target a variety of commodity computing engines in an HPC market where heterogeneous systems and computing environments are

the norm. HPC users are demanding compilers that enable a single high-level version of application source code which will re-compile and run with uniformly good performance on x86, POWER and ARM CPUs coupled to Tesla and Radeon accelerators, as well as on future self-hosted Xeon Phi systems. This talk will provide an overview of PGI's strategy and plans for delivering such compilers by leveraging the strengths of both proprietary and open source compiler technologies.

HPC Futures and Exascale

Chair: Lucas A. Wilson (Texas Advanced Computing Center)

1:30pm-3pm

Room: 292

Implementations of Hybrid Memory Cube for Networking and HPC Systems

Todd Farrell (Micron Technology, Inc.)

As we move into the next generation of high performance computing, the development and implementation of purpose built, intelligent memory subsystems are key to success. The past year has seen the acceleration of the adoption of Hybrid Memory Cube technology into high performance compute platforms, including Fujitsu's Post-FX10 system (unveiled at SC13) and Intel's second generation Xeon Phi architecture (Knights Landing). This presentation will provide an in-depth look at specific examples of HMC implementation in HPC and high performance networking systems. Additional detail will be shared on the positive impact to total cost of ownership at the system level, where HMC provides advantages in energy and board space efficiencies, as well as on overall system performance.

Exascale Ready RSC PetaStream Massively Parallel Supercomputer, RSC Tornado Cluster Solutions with World Records of Computing and Power Density are Proven by Customers in 2.5 PFLOPS Deployed Projects

Alexander Moskovsky (RSC)

RSC Group (www.rscgroup.ru), the leading Russian developer and system integrator of innovative energy efficient HPC and Data Center solutions, will talk at SC14 about RSC PetaStream - the revolutionary massively-parallel, ultra high density and energy efficient direct warm water cooled supercomputer solution which set world records of computing and power density: 1.2PFLOPS and 400kW per 1 m2 cabinet with 1024 Intel Xeon Phi 7120D. RSC PetaStream is focused to reach ExaScale levels of performance protecting software investments for HPC and Big Data applications. RSC liquid cooled systems are proven in operation since 2009. RSC PetaStream and RSC Tornado cluster architecture with Intel Xeon E5-2600 v3 (269TFLOPS and 100kW per cabinet) based solutions, RSC BasIS software stack

are deployed in large number of projects, including 1.1PFLOPS in Saint Petersburg State Polytechnic University.

HPC Workload Management on the Road to Exascale

Devin Jensen (Altair Engineering, Inc.)

From power consumption to affordability, the path to exascale isn't smooth. But the biggest barrier remains pure scalability – even if we can companies can supply the energy, footprint and cash needed for exascale machines, when will vendors be able to deliver the scale itself and the software needed to harness it for the promised speed and performance? Altair has been working toward exascale for many years, with the last two PBS Professional releases delivering major increases in speed and scalability. In this presentation we will describe the needs to date in preparing for exascale, provide an overview of the technology components needed to manage an exascale environment, and unveil Altair's latest results in scalability for our upcoming PBS Professional 13.0.

Software for HPC

Chair: Paul Domagala (Argonne National Laboratory)

1:30pm-3pm

Room: 291

3D in the Cloud - How NICE DCV 2014 Gets the Best from Virtual GPUs

Andrea Rodolico (NICE)

With the increasing concentration of large storage and HPC facilities, either in Public or Private Clouds, moving data to and from the user's desktop is getting more and more challenging. More and more often, companies move interactive jobs, including analysis, design and visualization tools, near the HPC resources. NICE DCV enables this remote mode of operation, moving pixels in place of large files, thanks to the GPU resources now increasingly available in different kinds of virtualized environments. In this talk we'll discuss the state of the art of GPUs and different virtualization options, explore the many remote-graphics-enabled options available in Private and Public Clouds, and see how DCV can leverage the latest features of these powerful resources to access virtually any kind of heavy duty visualization tool, both on Linux and on Windows, from any network, in full mobility.

Qluster: Cluster OS/Management with a Novel Communication Architecture

Roland Fehrenbacher (Q-Leap Networks GmbH)

We will give an introduction to Qluster, putting emphasis on the following: (a) The modular compute/storage node OS image technology and its advantages compared to conventional node provisioning. (b) The management framework QluMan

which is a fully integrated platform to configure and operate clusters at any scale. (c) The communication architecture of QluMan based on secure high-speed messaging.

High Performance Data Analytics: Experiences Porting the Apache Hama Graph Analytics Framework to an HPC Infiniband Connected Cluster

William Leinberger (General Dynamics)

Open source analytic frameworks provide access to Big Data in a productive and fault resilient way on scale-out commodity hardware systems. The objectives of High Performance Data Analytic systems are to maintain the framework productivity and improve performance for the data analyst. In order to achieve the performance available from High Performance Computing (HPC) technology, a framework must be recast from the distributed programming model common in the open source world to the parallel programming model used successfully in the HPC world by surgical replacement of key framework functions that leverage the strengths of HPC systems. We demonstrate this by porting the Apache Hama graph analytic framework to an HPC Infiniband Cluster. By replacing the distributed barrier class in the framework with a parallel HPC variant (prototyped in MPI), we achieved a performance increase of 37% on a real-world Community Detection application applied to a synthetic community rich graph.

HPC Futures and Exascale

Chair: Victor Eijkhout (University of Texas at Austin)

3:30pm-5pm

Room: 292

Fujitsu's New Supercomputer, Delivering the Next Step in Exascale Capability

Toshiyuki Shimizu (Fujitsu)

Fujitsu has been leading the HPC market for over 30 years, and today offers a comprehensive portfolio of computing products – high-end supercomputer PRIMEHPC series, x86-based PRIMERGY clusters, software, and solutions - to meet wide-ranging HPC requirements. At SC14 Fujitsu will demonstrate its sustained progress towards Exascale computing by highlighting its latest addition to the PRIMEHPC family.

Innovation and the Role of Instruction Set Architecture

John Hengeveld (ARM)

Most everyone agrees that the exascale era and beyond will require substantial innovation in system architecture. However, an interesting debate surrounds the role of specific Instruction Set Architecture as an attribute of system innovation. We will look beyond the standard meme of energy efficiency

and discuss how Instruction Set Architecture elements map to emerging issues in HPC. We'll also explore the necessary attributes for HPC ISA, HPC processors, and the system around them to address the workloads that are driving us into exascale. Finally, we'll look at how the ARM® ecosystem is driving innovation and how HPC researchers and developers can prepare for system deployments on ARMv8-based HPC systems.

Themes for HPC Exascale Storage

Nathan Rutman (Seagate Technology LLC)

None of today's storage systems can meet the extreme performance, availability, and consistency requirements of high-performance exascale computing. In this presentation we explore a set of themes to drive an exascale storage solution and describe how these themes may be embodied in a new, from-scratch solution.

Software for HPC

Chair: Jerome Vienne (University of Texas at Austin)

3:30pm-5pm

Room: 291

Forging Ahead in a Many-Core World: Does OpenMP Scale?

Mark O'Connor (Allinea Software)

Celebrating 10 years of high-quality development tools for science, Allinea takes a look ahead at the coming decade. HPC systems are nothing without the applications - and the trend towards mixed-mode MPI/OpenMP/Accelerator codes demands a stronger level of tool support than ever before. With lots of pictures and no absolutely bullet points this talk will look at the challenge of effectively developing and running these codes on both the current and the coming wave of systems.

Improving Productivity and Portability in HPC Using the ArrayFire Library

Shehzaan Mohammed (ArrayFire)

ArrayFire is an Open Source scientific computing library with a focus on portability and ease of use. ArrayFire's various compute backends and language bindings allow users to run highly performant code on CPUs, GPUs, FPGAs, accelerators, and other multi-core processors using a programming language of their choice.

In this presentation we will show how you can use the same high level code to get high performance across various systems and architectures. We will also demonstrate the modular nature of ArrayFire allows it to quickly adopt to newer architectures in the future.

Developing HPC Software in the Era of Multi-Level Parallelism

Chris Gottbrath, Wendy Hou (Rogue Wave Software, Inc.)

One of the essential challenges in both HPC and Big Data is the process of mapping some large computational problem to the available computational resources. HPC developers strive to create code that's efficient, scalable, reliable, maintainable and maps to compute resources that simultaneously require different kinds of parallelism: node parallelism, SMP parallelism across cores, offloading of work an accelerator and vector parallelism.

This talk will review these challenges and talk about how a selection of Rogue Wave's customers have tackled these challenges with by using IMSL analytics library or the TotalView debugger. It will show how these tools can be used at the workstation or laptop level, in departmental clusters or on dedicated high scale supercomputers. It will show how they help users who are taking advantage of Coarray Fortran 2008, C++ 2011, CUDA, OpenACC, MPI, OpenMP and the Intel Xeon Phi.

Thursday, November 20

Effective Application of HPC

Chair: Dane Skow (Self)

10:30am-12pm

Room: 291

HPC Without A Data Center

Devarajan Subramanian (Gompute Inc)

Find out how you can build better products with HPC, even without your own local data center!

SUSE Linux Enterprise 12 - Innovations for HPC

David Byte (SUSE)

SUSE Linux Enterprise 12 is bringing features that have specific benefits for HPC to the market. Topics of interest will include:
- Kernel Live Patching - Automatic NUMA Balancing - Support for bigger systems - NVMe support

An Introduction to the Center for Automata Processing

Kevin Skadron (University of Virginia)

The University of Virginia, with seed support from Micron Technology, Inc., has co-founded the Center for Automata Processing to catalyze an ecosystem around automata processing.

Micron's Automata Processor (AP) is a new, non-von-Neumann architecture capable of high-speed, symbolic search and analysis of complex, unstructured data streams. A highly scalable fabric of interconnected processing elements, the AP delivers massive, energy-efficient parallelism. The Center's goal is to coordinate academic, industry, and government researchers to advance the field of automata computing and train future data scientists and engineers. The Center conducts fundamental research on foundations and applications of automata computing as well as specific optimizations for the AP. This talk will describe several pilot projects that cover a variety of active research areas, including pattern mining, natural language processing, image analysis, and bioinformatics, with the goal of identifying applications that are well-suited for the AP and identifying important performance-optimization techniques.

Moving, Managing, and Storing Data

Chair: Wei Tang (Argonne National Laboratory)

10:30am-12pm

Room: 292

Building High Availability Storage in HPC Environment

Sergei Platonov (RAIDIX LLC)

Our presentation considers different approaches in building HA storage subsystems in HPC environment. We give definition to high availability and compare different alternatives to deploy parallel storage systems. The following options are opposed: shared disk cluster vs shared nothing cluster, replication vs erasure coding. We discuss evolution of RAID technology and describe development of RAIDIX core technologies, such as RAID 7.3, RAID 7.N, Cluster-in-a-Box for HPC Environment.

Building High Capacity Networks that Support Traffic Pattern Elasticity

Daniel Tardent (CALIENT Technologies)

Not only is data growing dramatically in data centers and high performance computing, but there's also much more traffic pattern variability and elasticity. For example, increased east-west data traffic together with Big Data applications such as Hadoop are resulting in more persistent data flows. These longer flows can saturate buffers, introducing unacceptable latency for shorter flows and seriously impairing the responsiveness of the network. This article will look at how a hybrid packet-optical switched network using the CALIENT S320 Optical Circuit Switch can provide a high-speed, fiber-optic overlay network to handle large flows, reducing network congestion dramatically.

Scalability Matters: Why Big Data Storage, Computing, and Analytical Tools Need Scalable Performance and Capacity

Joseph Landman (Scalable Informatics)

As data volumes continue to grow at exponential rates, the bandwidth and performance mismatch between edge-based storage and computing designs obviate their continued use. Simply slapping storage software onto inefficient hardware doesn't solve the problem, or even address it in any meaningful way. To be able to handle the big data problems of today and tomorrow, applications need to run where the data is, and the data needs to be able to be efficiently and rapidly accessed.

Effective Application of HPC

Chair: Si Liu (University of Texas at Austin)

1:30pm-3pm

Room: 291

The HPC Advisory Council Activities

Gilad Shainer, Brian Sparks (HPC Advisory Council)

The HPC Advisory Council's mission is to bridge the gap between high-performance computing (HPC) use and its potential, bring the beneficial capabilities of HPC to new users for better research, education, innovation and product manufacturing, bring users the expertise needed to operate HPC systems, provide application designers with the tools needed to enable parallel computing, and to strengthen the qualification and integration of HPC system products. The session will review the latest update on the HPC Center systems and resources, the latest publications on various applications, the collaborations with many of the world leading research labs and coming meetings and conferences.

HPC on Microsoft Azure

Robert Demb (Microsoft Corporation)

Microsoft provides a spectrum of HPC solutions that span the needs and requirements of users and developers.

Recently, Microsoft integrated these solutions into a framework called "BigCompute." This framework consists of APIs, services, and various compute infrastructures for expressing and structuring analytic and HPC workloads.

The emphasis of this framework is to run jobs, not manage clusters.

The most distinguishing feature of this framework is that it provides a holistic approach to designing and/or hosting existing and new applications in the Cloud.

This "better together" approach allows users and developers to get their results faster at a cheaper cost with a system that is optimized for their requirements.

We provide an overview of Microsoft's "better together" approach, and describe design points that range from virtualized RDMA interfaces that enable single digit latency, to more generalized SOA and declarative interfaces that support traditional MPI, parameter sweep, and general task farming workloads.

SysFera and ROMEO Make Large-Scale CFD Simulations Only 3 Clicks Away

Benjamin Depardon (SysFera), Arnaud Renard (University of Reims Champagne-Ardenne)

ROMEO, the world's 6th greenest supercomputer, is committed to helping SMBs make use of HPC resources by providing both computing resources and support. However, ease of use in today's HPC tools is lacking. This makes it difficult for non-computer scientists to run their simulations on supercomputers, and end-users are generally unaware of the Linux environment. To address these issues, SysFera offers SysFera-DS, a web-based portal that makes HPC easy by enabling remote interaction with and visualization of 3D applications, and regular HPC job submission.

The presentation will describe how SysFera-DS simplifies — through a simple web browser — the whole process of running large scale Computational Fluid Dynamics simulations on ROMEO's 2048 cores and 260 NVIDIA TESLA K20X for industrial use-cases. We will also present the simulation speedup and results, and the impact it had for Tech-Am Ingénierie, one of ROMEO's customers. Future evolutions and enhancements will be discussed.

Hardware and Architecture

Chair: Judicael A. Zounmevo (Argonne National Laboratory)

1:30pm-3pm

Room: 292

Introducing the "Brick", the New Modular, Accelerated, Water Cooled Aurora HPC Architecture

Giovanbattista Mattiussi, Paul Arts (Eurotech)

Eurotech presents their new Aurora Brick HPC architecture, which provides the basis for the new line of water cooled products presented at SC14. Eurotech designed this architecture to be very modular, dense and energy efficient. In this way, Eurotech aims to tackle some of the biggest challenges HPC is facing, offering a vertical special purpose HPC system that maximizes performance, minimizing energy consumption and space, with the same degree of flexibility of a generic purpose architecture. The basic concepts behind the Brick are modularity and boosting, so the possibility to considerably speed up both execution and storage, using multi-acceleration and fast storage, easily combined together by playing on the configuration of the HPC system through different PCIe modules. The Brick architecture supports a line of products that will use Intel, Nvidia and ARM (Applied Micro) as processors and co-processors in their first release, expected soon after SC14.

Using Basic Norms for Approximation of Node Associations

Bruce McCormick (CogniMem Technologies, Inc.)

A rapidly emerging theme in large-scale data mining is the use of models based on network theory (also known as graph theory), which allow abstract representation of associations (or “connections”) between entities (or “nodes”). There is growing consensus that basic norms, in spite of their simplicity often represent an acceptable approximation to more complex formulae, subject to appropriate data transformation. L1-norm is well-suited to represent correlation between time-series, and the Lsup-norm is relevant to situations where emphasis is given to mismatch even in a single parameter. An architecture, implemented in silicon and available from CogniMem Inc., represents an ideal platform for solving these problems, as it allows massively-parallel calculation of norms between a stored set of reference vectors and vectors that broadcast simultaneously to all processing units. Small clusters of parallel devices can easily exceed the performance of current multi-GHz CPUs, at a fraction of the power.

High-IQ Networks

JJ Jamison (Juniper Networks)

Software Defined Networking (SDN) will enable HPC network providers to create the kind of agile, high-IQ infrastructure needed to support the advancement of cutting edge scientific research.

Network systems must provide the necessary control points for use by the software and application layer. Software applications need to leverage network visibility and programmability to extract information and to convert it to knowledge and insight. With that, we can move from a world of limited visibility and no intelligence to a world of real-time visibility and smart systems.

This presentation will focus on how HPC networks can be designed and implemented in a way that enables the SDN system, hardware and software, to work seamlessly together and act automatically based on real-time network conditions. Both SDN solutions that automate and orchestrate the creation of highly scalable virtual networks and those that enable powerful and flexible traffic engineering will be discussed.

Effective Application of HPC

Chair: Richard Evans (Texas Advanced Computing Center)

3:30pm-5pm

Room: 291

Cray Matter: Brilliance-Boosting Technologies

Barry Bolding (Cray Inc.)

Why does Cray matter? Because ever since our founder Seymour Cray introduced the world’s first supercomputer, our company has recognized that in order for the world’s “brains” to make their remarkable breakthroughs in science and engineering, they need powerful HPC tools to turn those great ideas into brilliant discoveries. Today more than ever, Cray is arming these scientific brains with a new arsenal of computing, storage and analytics solutions to solve the most complex problems and get to the solution faster. Join us to hear Steve Scott talk about why Cray — and HPC — matter.

Hybrid and Public Cloud Optimized for Technical and High Performance Computing

Louise Westoby (Platform Computing), Chris Porter (IBM Corporation)

While cloud computing presents many opportunities for organizations to improve efficiency and realize bottom-line business advantages, there are also barriers to adopting public cloud for HPC— such as security, licensing, pricing and performance. The need for scale and compute power is without question, but there is debate about whether cloud infrastructure providers can handle the needs of compute-intensive applications, and under what circumstances. This session will review what to look for when evaluating which infrastructure as a service (IaaS) provider is best for your needs. We will introduce you to the IBM® Platform Computing™ Cloud Service running on SoftLayer®, which delivers high performance application-ready clusters in the cloud for organizations that need to quickly and economically add computing and storage capacity. You can deploy analytics and technical computing workloads on hybrid and secure stand-alone clusters in the cloud, leveraging dedicated, bare metal infrastructure and InfiniBand for optimal performance and security.

Big Data Changes Everything: Reinventing HPC

Dave Turek, Ruud Haring (IBM Corporation)

HPC must evolve in two fundamental ways: first, target complete workflows of critical interest to the organization; second, explicitly accommodate the impact of Big Data on system/solution designs. Today, many “HPC solutions” are nothing more than optimizations of important but narrowly defined algorithms embedded in much more involved workflows. The other elements of the workflow, including issues of data management and manipulation coupled with associated analytics, are often ignored. As a consequence, the value of HPC to the end user is less than it could be. As we enter the second half of this decade, IBM intends to reconcile these issues: take workflows that matter and apply to them all the HPC and analytic tools required to provide the greatest insight possible in the shortest amount of time. This session will talk to requirements and announcements that will successfully deliver innovations for the next generation of analytics and technical computing.

Software for HPC

Chair: Si Liu (University of Texas at Austin)

3:30pm-5pm

Room: 292

Accelerate Insights and Tame Complexities with Moab

Paul Anderson (Adaptive Computing)

Managing and analyzing unprecedented amounts of diverse data presents a complex problem for companies today — even with the availability of sophisticated modeling, simulation and process tools. In this session, Adaptive Computing will share their experiences managing the world’s largest supercomputing centers in the academic, government and commercial sectors and demonstrate new features such as HPC cloud bursting, an administrator portal and higher performance and scale capabilities.

Adaptive will discuss new ways to:

- Unify data center resources in order to break down siloed environments and create an adaptive ecosystem
- Optimize the analysis process by automating workflows, maximizing utilization and improving resource performance
- Guarantee services to the business to accelerate insights!

Armed with the insights from this session, attendees will be able to reduce the complexity surrounding intense simulations and big data analysis to accelerate insights more rapidly, accurately and cost-effectively.

Bright Cluster Manager: A Managed Solution for Data-Aware HPC On Premise and In The Cloud

Martijn de Vries (Bright Computing, Inc.)

Effective HPC demands locality between application data and compute resources. Therefore, use cases involving the extension of on-premise HPC resources into the cloud must address this locality requirement. Bright Cluster Manager ensures data-aware scheduling of HPC workloads in Amazon Web Services (AWS) by coordinating the instantiation of compute resources (in the AWS cloud) with the local availability of data through native support for Amazon S3 and Glacier. Bright’s CLI or GUI includes supports for Amazon VPC setups as well as use of cloud-based GPU instances in hybrid-architecture compute resources. Bright originally introduced the extension of on-premise HPC resources into AWS in 2011. In this presentation, we will preview our latest product enhancements that address the locality demands of data and compute from a more-holistic storage perspective for HPC as well as Big Data Analytics.

SaltStack for Fast and Scalable Automation of HPC Distributed Systems

Jeff Porcaro (SaltStack)

The bigger the data, the bigger the requirement for more computing power. If we can automate the build out and administration of the infrastructure that makes super computing possible, the more value we can extract out of data to deliver the real benefits of high performance computing. SaltStack is systems and configuration management software used to define the software-defined data center and to manage all the things including compute, storage, and the network. We know data centers aren’t built like they were even five years ago, and today’s infrastructure requirements require more scale, security, stability, flexibility, and open standards. This presentation will show how SaltStack is used by organizations like the Center for Disease Control, Technicolor, LSU and Argonne National Laboratory to build a modern, software-defined infrastructure in minutes. The SC environment you always wished you had, shouldn’t be too difficult to get.

HPC Impact Showcase/ Emerging Technologies

HPC Impact Showcase/ Emerging Technologies

HPC Impact Showcase

The HPC Impact Showcase reveals real-world HPC applications via presentations in a theater setting on the exhibit floor.

The Showcase is designed to introduce attendees to the many ways HPC is shaping our world through testimonials from companies, large and small, not directly affiliated with an exhibitor. Their stories relate real-world experience of what it took to embrace HPC to better compete and succeed in their line of endeavor.

Whether you are new to HPC or a long-time professional, you are sure to see something new and exciting in the HPC Impact Showcase. Presentations will be framed for a non-expert audience interested in technology and will discuss how the use of HPC has resulted in design, engineering, or manufacturing innovations (and will not include marketing or sales elements!).

Emerging Technologies

SC14 provides a showcase for high-risk, high-reward hardware and software technologies that may significantly change the world of HPC in the next ten years. For example, technologies like reconfigurable computing, new SoC designs, alternative programming systems, and novel cooling techniques may offer near-term benefits, while new device technologies, like carbon nanotubes, non-volatile memory, quantum computing, and chip-level optical interconnects offer potentially paradigm-changing benefits over the long term.

HPC Impact Showcase/ Emerging Technologies

Monday, November 17

Emerging Technologies Exhibits

*Chairs: Satoshi Matsuoka (Tokyo Institute of Technology),
Jeffrey Vetter (Oak Ridge Laboratory)*

7pm-9pm

Room: Show Floor – 233 Exhibits

DEEP Collective Offload

Florentino Sainz, Vicenç Beltran, Jesús Labarta (Barcelona Supercomputing Center), Carsten Clauss, Thomas Moschny (ParTec Cluster Competence Center GmbH), Norbert Eicker (Juelich Research Center)

Exascale performance requires a level of energy efficiency only achievable with specialized hardware. Hence, the trend is clearly going towards heterogeneous architectures. However, in order to fully exploit the advantages of such heterogeneous hardware, there is a need for appropriate programming models that do not increase the complexity of the applications. In the DEEP project, we have extended the OmpSs programming model to offload MPI kernels dynamically by utilizing MPI's spawning feature, here namely provided by ParaStation MPI as the underlying infrastructure. That way, offloading to specialized hardware gets dramatically simplified while hiding the low-level and error prone MPI spawn calls. These facts together with its portability make this concept definitely interesting for a broader audience of the HPC community.

Persistent In-Memory Computing with Emerging Non-Volatile Memory

Qingsong Wei, Jun Yang, Cheng Chen, Mingdi Xue, Chundong Wang, Renuga Kanagavelu (Data Storage Institute)

Emerging non-volatile Memory (NVM) is changing the way we think about memory and storage because they provide DRAM-like performance, disk-like high density and persistence. Current system is designed and optimized for traditional DRAM-Disk hierarchy. When memory is replaced with NVM, memory is becoming persistent. The inefficiency of current architecture and system software significantly limits the potential of NVM. The objective of this project is to redesign the system and address architectural challenges and develop a suite of techniques which enable persistent in-memory computing with NVM for high performance computing. The major works we have achieved are NVM-accelerated File System and Consistent and Durable In-memory Key-value Store. NVM-accelerated File System explores NVM's byte-addressability and persistence to improve metadata latency and reduce metadata I/O. Consistent and Durable In-memory Key-value Store achieves high performance through designing NVM-based indexing

data structure with low overhead of maintaining data consistency.

InfiniCortex: Concurrent Supercomputing Across the Globe Utilizing Trans-Continental InfiniBand and Galaxy of Supercomputers

*Marek Michalewicz, Yuefan Deng, Tin Wee TAN, Yves Poppe (A*Star Computational Resource Center), Scott Klasky (Oak Ridge National Laboratory), David Southwell (Obsidian Strategics)*

We propose to merge four concepts integrated for the first time together to realise InfiniCortex demonstration: (1) High bandwidth intercontinental connectivity between Asia and the USA; (2) InfiniBand technology over trans-continental distances using Obsidian range extenders; (3) Connecting separate InfiniBand sub-nets with different net topologies to create a single computational resource: Galaxy of Supercomputers; (4) Running workflows and applications on such a distributed computational infrastructure, especially using ADIOS I/O framework.

Our demo set up of InfiniCortex will comprise of computing resources residing in Singapore (A*STAR), Japan (Tokyo Institute of Technology), Australia (National Computational Infrastructure) and in the United States (University of Tennessee), as well as on the floor of the Exhibition hall, at the A*CRC booth. These computers will be connected using long-distanced InfiniBand links among these five sites enabled by Obsidian Strategics' SDR Longbow E100 encrypting IB range extenders and QDR Crossbow R400 InfiniBand routers.

Tools and Techniques for Runtime Optimization of Large Scale Systems

Martin Schulz, Abhinav Bhatale, Peer-Timo Bremer, Todd Gamblin, Kathryn Mohror, Barry Rountree (Lawrence Livermore National Laboratory)

Future HPC systems will exhibit a new level of complexity, in terms of their underlying architectures and system software. At the same time, the complexity of applications will rise sharply, to both enable new science and exploit the new hardware features. Users will expect a new generation of tools and runtime support techniques that help address these challenges, work seamlessly with current and new programming models, scale to the full size of the machine, and address new challenges in resilience and power-aware computing.

We present a wide range of tool and runtime system research efforts targeting performance analysis, performance visualization, debugging and correctness tools, power-aware runtimes, and advanced checkpointing. These efforts are supported by

activities on tool and runtime infrastructures that enable us to work with modular tool designs and support rapid tool prototyping. We showcase our approaches to users, system designers and other tool developers.

Can Containers Enable the Convergence of HPC, Big Data and Cloud?

Seetharami Seelam, I-Hsin Chung (IBM Research), Guangpu Li, Cherung Lee (National Tsinghua University)

Operating system level containers or containers for short are operating-system level virtualization technology that enable execution of multiple isolated operating systems on the host operating system. In this emerging technology project we will demonstrate how containers could be used both on traditional HPC systems and on cloud platforms to execute traditional HPC and Big Data and Enterprise applications on the shared resources.

All-to-All, Low Latency, and High-Throughput Optical Interconnect Demonstrator for Scalable Supercomputing

Zheng Cao, Roberto Proietti, S. J. Ben Yoo (University of California, Davis)

This project demonstrates five key elements for future high-performance computing: (1) high-throughput; (2) low-latency; (3) energy-efficiency; (4) scalability; and (5) high-radix, all-to-all connectivity. We exploit optical parallelism and wavelength routing capability of arrayed wavelength grating routers (AWGRs) and collapse the entire network topology to a flattened and single-hop interconnection topology. The AWGR is a passive wavelength routing component that enables fully connected all-to-all parallel interconnection without contention across huge (>20 THz) optical bandwidths. This project demonstrates wavelength routed optical interconnects that (1) scale to ~ million compute nodes and Exabyte/s bisection-bandwidth, (2) support all-to-all connectivity without contention or arbitration with the radix count 512 and beyond, (3) exploit ~ pJ/bit communication energy efficiency, (4) achieves 100% throughput at the rack and 97% throughput in a cluster at 100% data injection rate on every transmitter port in the system, and (5) demonstrates a chip-scale AWGR switch implementation using silicon photonics.

The AweSim Appkit: A Platform for Developing Web-Based HPC Applications

Alan Chalker, Dave Hudak (Ohio Supercomputer Center)

The Ohio Supercomputer Center's AweSim program is working to lower the barriers to entry for HPC based modeling and simulation. This initiative is designed to simplify the use of advanced simulation-driven design with integrated point-and-click apps to dramatically lower the cost of such simulations.

AweSim uses four key technologies: (1) advanced models for manufacturing processes (2) a simple web-based user interface, (3) cloud-based computing resources and (4) computational solver software. The App Kit is a multilayer architecture to tie these components together, consisting of a web access layer, a web application layer and a system interface layer. The web access layer provides integrated single sign-on authentication with the App Store and industry-standard web proxy services. The web application layer is built on industry-standard technologies and contains reusable app templates. Finally, the system interface layer provides direct access to local resources such as databases, file systems and HPC job queues.

Fast-Reconfigurable Optical Networks for Data-Intensive Computing

Benjamin G. Lee, Clint L. Schow, Marc A. Taubenblatt, Fabrizio Petrini (IBM Corporation)

Data-intensive computing increasingly involves operations at the scale of an entire system, requiring quick and efficient processing of massive datasets. We present a network architecture, comprised of optical-switch and burst-mode transceiver technology, designed to support demanding graph algorithms in distributed-memory systems. The fabric is designed to deliver up to 10 terabytes per second of node bandwidth and predictable performance under heavy load with latencies under a microsecond. This vision for a unique optical switch provides a path to overcoming the limitations of electrical switch networks. To make this a reality, we are developing nanosecond-scale photonic switches and low-power 25-Gb/s WDM transceivers with burst-mode clock and data recovery for link retraining in tens of nanoseconds. Network simulations predict graph performance on par with today's leading supercomputers, and orders of magnitude improvement in power efficiency and footprint would allow the system to fit within a few racks.

Chapel Hierarchical Locales: Adaptable Portability for Exascale Node Architectures

Greg Titus, Sung-Eun Choi, Bradford L. Chamberlain (Cray Inc.)

As HPC approaches exascale computing, processor and node architectures are diverging rapidly, between distinct vendors and even across generations from a given vendor. This challenges programming models that cannot expose intra-node locality or rapidly adapt to novel processor architectures. It also challenges end-users by requiring them to modify their programs when moving between systems.

Our emerging technology exhibit introduces Chapel's hierarchical locales—a programming language abstraction that facilitates rapid adaptation to such architectural changes. Hierarchical locales differ from previous approaches in that users can provide their own architectural models to the compiler,

providing the mapping from high-level parallel algorithms using standard Chapel concepts down to target systems. Moreover, hierarchical locales impose fewer semantic restrictions than previous approaches.

In this exhibit, we introduce the notion of hierarchical locales, explain how they are deployed and used in the Chapel implementation today, and report on our support for NUMA nodes and Intel MIC.

A New Node Design for HPC

Roger Ronald, Kevin Moran (System Fabric Works), Thomas Sohmers, Kurt Keville, Paul Sebexen (Rex Computing)

Current trends in processor technology have led to bottlenecks in memory latency, data movement, and storage. A need exists to accelerate practical exascale computing by focusing on power efficiency, simplicity of component design, and hierarchically designed systems. The only way to build a truly high performance computing system capable of meeting exascale demands and beyond is to design the simplest and thus most efficient processing unit, without sacrificing functionality or scalability.

System Fabric Works and REX Computing have been jointly developing a full-node architecture centered around a new processor approach optimized for HPC. We propose to build a grid of many (potentially thousands of) Multiple Instruction, Multiple Data (MIMD) cores to execute instructions independently, drastically reducing the demands on memory bandwidth and thus greatly increasing the efficiency at which typical HPC workloads may be performed.

The ASTRON - IBM 64bit μ Server for SKA

Ronald Luijten (IBM Zurich Research Laboratory)

For the IBM-ASTRON DOME μ Server project, we are currently finishing two compute node boards in memory DIMM-like form factor. The first is based on a 4 core 2.2 GHz and the second on a 12 core / 24 thread 1.8 GHz SoC. Both SoCs are available from Freescale employing 64-bit Power ISA. These SoCs meet our μ Server definition: integrate the entire server motherboard into a single microchip, except DRAM, NOR-boot and power conversion circuitry. Our focus is on density yet also on compute, memory and I/O BW balance. Our innovative hot-water based cooling infrastructure also supplies the electrical power to our compute node board, which enables us to put up to 128 compute nodes in a single 2U rack-unit. This unit will provide 3000 HW threads, 6 TB of DRAM and 32 Ethernet interfaces at 40Gbps. For SC14 we would demo single compute node air-cooled systems running Fedora20 and DB2.

FPGA Development Tools and HMC/FPGA Architectures

Steve Wallach (Convey Computer Corporation)

Convey has developed an OpenMP compiler that accepts standard OpenMP syntax and outputs a thread level intermediate language – HT (Hyper_Threading Tools). The HT tool sets, further compiles the intermediate language to Verilog.

The system developed can be used to create “FPGA personalities” for Convey developed hardware platforms. These platforms range from PCI cards, complete compute nodes, and systems that use HMC memory architectures

Blamange: Fat Binaries on a Diet

Kurt Keville (Massachusetts Institute of Technology), Anthony Skjellum (Auburn University), Mitch Williams (Sandia National Laboratories)

The latest generation of SoCs present new opportunities for performance and reliability by integrating multiple architectures in silica. At SC14 we intend to demo a cluster that takes advantage of a novel functionality within CUDA 6.0 that allows quick pivots and context switches within shared memory. This demo will show new strategies for hardware-level cyber-defense. The takeaway from this presentation is how the application is facilitated through fat binaries, with its broad implications to HPC.

DEEP-ER I/O: An Exascale I/O Framework

Sven Breuner (Fraunhofer Institute for Industrial Mathematics ITWM), Giuseppe Congiu (Xyratex), Norbert Eicker, Wolfgang Frings (Juelich Research Center), Sai Narasimhamurthy (Xyratex), Kay Thust (Juelich Research Center)

This work describes DEEP-ER I/O, a highly scalable I/O solution potentially suitable for exascale platforms. The work addresses the I/O requirements of exascale applications in the DEEP-ER project, which is one of the European exascale computing initiatives funded by the EU. DEEP-ER I/O exploits the complimentary I/O optimization and scalability aspects of SIONlib, which primarily addresses problems due to a large numbers of small files, Exascale10 based ExCol solution, which overcomes issues due to small I/O, and the BeeGFS file system and its extensions, which intelligently uses NVRAM tiers in the I/O stack. The solution addresses performance, power, reliability and cost issues associated with scaling I/O in HPC.

Monday, November 17

Gala HPC Impact Showcase

Chair: David Martin (Argonne National Laboratory)

7:15pm-8:30pm

Room: Show Floor - 233 Theater

Introduction to the HPC Impact Showcase

David E. Martin (Argonne National Laboratory), David Halstead (National Radio Astronomy Observatory)

Overview and highlights from of the exciting presentations to be given by industry leaders throughout the week. Welcome to the HPC Impact Showcase!

The Exascale Effect: the Benefits of Supercomputing Investment for U.S. Industry

Cynthia R. McIntyre, Chris Mustain (Council on Competitiveness)

The Council on Competitiveness has worked for over a decade to analyze, promote and strengthen the America's use of advanced computing for economic competitiveness. The Council brings together top high performance computing leaders in industry, academia, government and the national laboratories. We know from experience that to out-compete is to out-compute.

Supported by a grant from the U.S. Department of Energy, the Council recently issued the report *Solve – the Exascale Effect: the benefits of Supercomputing Investment for U.S. Industry*. The report takes a fresh look at: (1) the value of HPC to U.S. industry, (2) key actions that would enable companies to leverage advanced computing more effectively, and (3) how American industry benefits directly and indirectly from government investment at the leading edge of HPC.

The leaders of the Council's HPC Initiative, Senior Vice President Cynthia McIntyre and Vice President Chris Mustain, will share the findings of the report.

International Data Corporation Study for DOE: Evaluating HPC ROI

Steve Conway (International Data Corporation)

Across the world, government funding for large supercomputers is undergoing a multi-year shift. The historical heavy tilt toward national security and advanced science/engineering increasingly needs to be counterbalanced by arguments for returns-on-investment (ROI). The U.S. Department of Energy's (DOE) Office of Science and National Nuclear Security Administration recently awarded International Data Corporation (IDC) a three-year grant to conduct a full study of returns on investments (ROI) in high performance computing (HPC). The full

study follows IDC's successful completion of a 2013 pilot study on this topic for DOE. IDC's Steve Conway will discuss the goals and methodology for the three-year, full-out study.

Tuesday, November 18

Emerging Technologies Exhibits

10am-6pm

Room: Show Floor – 233 Exhibits

HPC Impact Showcase Presentations

Chair: David Martin (Argonne National Laboratory)

10:20am-2:20pm

Room: Show Floor – 233 Theater

HPC Productivity - Tearing Down the Barriers to HPC for CAE Simulation

Roger Rintala, Steve Legensky (Intelligent Light)

Engineering users are well aware of the familiar hurdles of exploiting HPC: Generating data faster than it can be evaluated, software scalability, licensing policies that limit production deployment and the need to keep internal resources working on tasks that are "on mission" for the organization.

Learn how Corvid Technologies, an engineering services and software provider, and longtime user of HPC, is delivering high quality, CAE-driven engineering for applications where accuracy and timeliness are critical. They are using the DoE's ultra-scalable, open source VisIt software with their highly scalable Velodyne solver to overcome the obstacles to HPC faced by many in the engineering community.

The use of VisIt brought with it some new challenges common with open source software. Corvid turned to Intelligent Light to provide them with professional support and some custom VisIt development that helped them complete their workflow and ensure that VisIt is meeting their HPC-enabled engineering needs. This development work produced multiple improvements to the VisIt codebase that have been contributed to the VisIt repository so they may benefit other VisIt users.

Intelligent Light, the US Department of Energy and users like Corvid are tearing down the barriers to the growth of HPC and improving the ROI on HPC-enabled CAE. They are doing so in ways that protect the choices of users and the investments they make in their simulations.

Impact of HPC on Nuclear Energy: Advanced Simulations for the Westinghouse AP1000® Reactor

Fausto Franceschini (Westinghouse Electric Corporation)

Westinghouse Electric Corporation, a leading supplier of nuclear plant products and technologies, performed advanced simulations on the Oak Ridge Leadership Computing Facility that reproduce with unprecedented fidelity the conditions occurring during the startup of the AP1000 advanced nuclear reactor. These simulations were performed using advanced modeling and simulation capabilities which are developed by the Consortium for Advanced Simulation of Light Water Reactors (CASL), a U.S. Department of Energy (DOE) Energy Innovation Hub led by the Oak Ridge National Laboratory. The AP1000 reactor is an advanced reactor design with enhanced passive safety and improved operational performance, with eight units currently being built, four in China and four in the U.S.

The high-fidelity simulation toolkit used for this effort corroborated a complete understanding of the behavior for this advanced reactor design. The novel simulations performed have been accomplished using a 60 million core-hour allocation Early Science award on the Oak Ridge Leadership Computing Facility TITAN. Use of the CASL advanced nuclear methods and software techniques, properly designed to take advantage of these vast computational resources, has been a key factor to this successful endeavor. Prior to this project, results of comparable fidelity were unavailable due to the impractical computational time resulting from limitations in codes scalabilities or computational resources available.

This endeavor, pursued through a collaborative effort of U.S. national labs and private industry, typifies the benefits that HPC can bring to the U.S. public by further ensuring the viability and safety of advanced energy sources.

HPC and Open Source

Omar Kebbeh (Deere & Company)

Open Source software has been around for generations. Some were born at universities and died there, others were born in national labs and made it out in terms of ISV codes, and others were born at universities and stayed as open where many scientists keep them alive.

This talk will address some of the issues faced by the industry when it comes to Engineering Open Source Codes. Both availability and challenges will be discussed and posed as an "Open Question" to the participants.

HPC Impact on Developing Cleaner and Safer Aero-Engines

Jin Yan (General Electric)

An aeroengine is an essential part of aviation industry. Over the past 40 years, the fuel efficiency and emissions of the aeroengine have improved by approximately 75% largely due to the continual increase of the combustor temperatures. Increasing the combustor exit temperature carries with it many challenges to the design. General Electric (GE) has continuously invested in High Performance Computing to address these challenges. Today, researchers and designers leverage state of the art aerodynamics and multiphysics tools to understand the finest flow scales and how small changes impact the overall performance of the aeroengines. This talk will provide an overview of HPC at GE today, and outlook on the future of large scale simulations in the aviation industry.

Scalable HPC Software – "It Takes a Village"

Mike Trutt (Northrop Grumman)

In today's competitive world, HPC requires software that can scale from just a few cores up to thousands. This requires a diversity of talent and skills that promote the creation of scalable software that maximizes utilization of available resources. There are several steps to creating scalable HPC software. Whether the application requires thousands of cores (like Computational Fluid Dynamics (CFD)) or is limited to a few dozen cores (such as a real-time signal processing platform), whether the hardware is commodity or custom, the discipline and principles for what makes software scale remain the same. Northrop Grumman's approach to producing scalable HPC software begins with modeling and simulation activities that feed software design requirements. The software in turn is designed for scalability, including such factors as middleware, scientific libraries, and benchmarking. Agile management can be key, as such approaches allow for accurate planning and course correction. Through cross-discipline interaction and training across the software lifecycle, scalable HPC software can enable today's endeavors in an increasingly competitive world.

Emerging Technologies Presentations

4:30pm-6pm

Room: Show Floor - 233 Theater

Chapel Hierarchical Locales: Adaptable Portability for Exascale Node Architectures

Greg Titus, Sung-Eun Choi, Brad Chamberlain (Cray Inc.)

As HPC approaches exascale computing, processor and node architectures are diverging rapidly, between distinct vendors and even across generations from a given vendor. This challenges programming models that cannot expose intra-node locality or rapidly adapt to novel processor architectures. It also challenges end-users by requiring them to modify their programs when moving between systems.

Our emerging technology exhibit introduces Chapel's hierarchical locales—a programming language abstraction that facilitates rapid adaptation to such architectural changes. Hierarchical locales differ from previous approaches in that users can provide their own architectural models to the compiler, providing the mapping from high-level parallel algorithms using standard Chapel concepts down to target systems. Moreover, hierarchical locales impose fewer semantic restrictions than previous approaches.

In this exhibit, we introduce the notion of hierarchical locales, explain how they are deployed and used in the Chapel implementation today, and report on our support for NUMA nodes and Intel MIC.

InfiniCortex: Concurrent Supercomputing Across the Globe Utilizing Trans-Continental InfiniBand and Galaxy of Supercomputers

*Marek Michalewicz, Yuefan Deng, Tin Wee TAN, Yves Poppe (A*STAR Computational Resource Centre), Scott Klasky (Oak Ridge National Laboratory), David Southwell (Obsidian Strategies)*

We propose to merge four concepts integrated for the first time together to realise InfiniCortex demonstration:

i) High bandwidth intercontinental connectivity between Asia and the USA; ii) InfiniBand technology over trans-continental distances using Obsidian range extenders; iii) Connecting separate InfiniBand sub-nets with different net topologies to create a single computational resource: Galaxy of Supercomputers iv) Running workflows and applications on such a distributed computational infrastructure, especially using ADIOS I/O framework.

Our demo set up of InfiniCortex will comprise of computing resources residing in Singapore (A*STAR), Japan (Tokyo Institute of Technology), Australia (National Computational Infrastructure) and in the United States (University of Tennessee) as well as on the floor of SC14 Exhibition hall, at the A*CRC booth. These computers will be connected using long-distanced InfiniBand links among these five sites enabled by Obsidian Strategies' SDR Longbow E100 encrypting IB range extenders and QDR Crossbow R400 InfiniBand routers.

All-to-All, Low Latency, and High-Throughput Optical Interconnect Demonstrator for Scalable Supercomputing

Zheng Cao, Roberto Proietti (University of California, Davis), S. J. Ben Yoo (University of California, Davis)

This project demonstrates five key elements for future high-performance computing: (1) high-throughput, (2) low-latency, (3) energy-efficiency, (4) scalability, and (5) high-radix, all-to-all connectivity. We exploit optical parallelism and wavelength routing capability of arrayed wavelength grating routers (AWGRs) and collapse the entire network topology to a flattened and single-hop interconnection topology. The AWGR is a passive wavelength routing component that enables fully connected all-to-all parallel interconnection without contention across huge (>20 THz) optical bandwidths. This project demonstrates wavelength routed optical interconnects that (1) scale to ~ million compute nodes and Exabyte/s bisection-bandwidth, (2) support all-to-all connectivity without contention or arbitration with the radix count 512 and beyond, (3) exploit ~ pJ/bit communication energy efficiency, (4) achieves 100% throughput at the rack and 97% throughput in a cluster at 100% data injection rate on every transmitter port in the system, and (5) demonstrates a chip-scale AWGR switch implementation using silicon photonics.

HPC Impact Showcase Presentations

Chair: David Martin (Argonne National Laboratory)

10:20am-2:20pm

Room: Show Floor – 233 Theater

HPC: A Matter of Life or Death

Terry Boyd (Centers for Disease Control and Prevention)

This session will use the example of a fictitious “zombie virus” outbreak to demonstrate the multiple applications of HPC to public health. These include: 1. Big Data Analysis to detect outbreaks 2. Spatial modeling and modeling of potential outbreaks to develop emergency response plans 3. Genomic evaluation of suspected disease outbreaks 4. Drug manufacturing and supply modeling 5. Contact tracing and response modeling to evaluate current activities and improve response activities 6. Post-event analysis to prepare for the next epidemic.

Though this session will use a fictitious zombie outbreak as the basis for discussion in order to increase interest in the session, the uses described will be applicable to real life outbreaks such as Ebola and H1N1.

Wednesday, November 19

Emerging Technologies Exhibits

10am-6pm

Room: Show Floor – 233 Exhibits

HPC Impact Showcase Presentations

David Halstead (National Radio Astronomy Observatory)

10:20am-2:20pm

Room: Show Floor – 233 Theater

High Performance Computing and Neuromorphic Computing

Mark Barnell (Air Force Research Laboratory)

The Air Force Research Laboratory Information Directorate Advanced Computing Division (AFRL/RIT) has focused research efforts that explore how to best model the human visual cortex. In 2008 it became apparent that parallel applications and large distributing computing clusters (HPC) were the way forward to explore new methods. In 2009 at our High Performance Computing Affiliated Resource Center (HPC-ARC) we designed and built a large scale interactive computing cluster. CONDOR, the largest interactive Cell cluster in the world, integrated heterogeneous processors of IBM Cell Broadband Engine multicore CPUs, nVidia GPGPUs and powerful Intel x86 server nodes in a 10GbE Star Hub network and 20Gb/s Infini-band Mesh, with a combined capability of 500 Teraflops. Applications developed and running on CONDOR include: neuromorphic computing applications, video synthetic aperture radar (SAR) backprojection and Autonomous Sensing Framework (ASF). This presentation will show progress on performance optimization using the heterogeneous clusters and how neuromorphic architectures are advancing the capabilities for the autonomous systems within the Air Force.

High-Order Methods for Seismic Imaging on Many-core Architectures

Amik St-Cyr (Shell Oil Company)

The lion's share of Shell's global HPC capacity is consumed by geophysical seismic imaging. Legacy algorithms and software must be replaced with fundamentally different ones that scale to 1000s of possibly heterogeneous-cores. Geophysical Reverse Time Migration is an example of a wave based imaging algorithm example. In this talk, we present how we're adapting our algorithms to tackle the manycore phenomena and how HPC impacts our business.

Using Next Gen Sequencing and HPC to Change Lives

Shane Corder (Children's Mercy Hospital)

After nearly 30 years, the Human genome is now being used in the clinical space. The researchers and engineers of The Center for Pediatric Genomic Medicine (CPGM) at Children's Mercy

Hospital in Kansas City, Missouri use HPC to decipher the infant's genome to rapidly look for genetic markers of disease to make a dramatic impact on the lives of children. CPGM is using HPC to push the turnaround time for this life saving application to just 50 hours.

Real-Time Complex Event Processing in DSPs

Ryan Quick (PayPal), Arno Kolster (PayPal)

Traditionally, complex event processing utilizes text search and augmented off-line analytics to provide insight in near real-time. PayPal's pioneering work leverages the true real-time capabilities of digital signal processors, event ontology, and encoding technologies to provide real-time pattern recognition and anomaly detection. Their novel approach delivers not only true real-time performance, but does so at a fraction of the power cost generally associated with low latency analytics. Partnering with HP and Texas Instruments, PayPal is delivering HPC solutions which not only capitalize on the power of offload compute, but also pave the way for new methodologies as we march towards exascale.

Lockheed-Martin's Special Multilevel and Overlapping Tier Support for the DoD HPC Centers

James C. Ianni (Lockheed Martin Corporation and Army Research Laboratory), Ralph McEldowney (Department of Defense HPC Modernization Program)

This presentation will discuss the benefit to the U.S. Warfighter provided by Lockheed Martin's multiple levels of HPC resource support for the various Department of Defense (DoD) Supercomputing Resource Centers (DSRCs). The DoD HPC Modernization Program supports the HPC needs of the DoD Research, Development, Test, and Evaluation community through several of the DCRCs: the Army Research Laboratory DSRC, the Navy DSRC, the Air Force Research Laboratory DSRC, and the Engineer Research and Development Center DSRC. Each of these centers has several state-of-the-art supercomputers upon which DoD scientists and engineers run high-fidelity simulations in fluid dynamics, structural mechanics, computational chemistry and biology, weather prediction, electromagnetics and other related disciplines. These simulations ultimately provide a great benefit to the U.S. Warfighter. Consequently, this involves a synergistic relationship between account creation, software utilization (commercial, open-source and custom) and optimization (compiler, code and/or theory utilization), hardware utilization and maintenance and user interaction. These synergies are enabled by cultivating a means of communication between the several tier groups within and across the individual Centers. This synergy is very important for HPC user interaction, as user requests can (and often do) span many levels of experience. These levels include getting an account and connecting to an HPC resource, compiling/submitting/script writing, and determining how to calculate spectra for an excited state, high-spin, charged,

f-element in various solvents. This presentation will describe the experiences of incorporating and actively utilizing all of the above-mentioned HPC support, as well as the ultimate impact HPC has had on the DoD.

Digital Consumer Products: Surfactant Properties from Molecular Simulations

Peter H. Koenig (Procter & Gamble)

Wormlike Micelles (WLMs) provide the basis for structure and rheology of many consumer products. The composition, including concentration of surfactants and level of additives such as perfumes and salt, critically controls the structure and rheological properties of WLM formulations. An extensive body of research is dedicated to the experimental characterization and modeling of properties of WLMs. However, the link to the molecular composition and micellar scale has only received limited attention, limiting rational design of WLM formulations. We are developing and refining methods to predict properties of micelles including cross-sectional geometry and composition, persistence length (relating to the bending stiffness) and scission energy using molecular dynamics simulations. This presentation will give an overview of the modeling program with participants from Procter and Gamble, U. Michigan, XSEDE, Oak Ridge National Labs and U. Cincinnati.

The ASTRON - IBM 64bit μ Server for SKA

Ronald Luijten (IBM Corporation)

For the IBM-ASTRON DOME μ Server project, we are currently finishing two compute node boards in memory DIMM-like form factor. The first is based on a 4 core 2.2 GHz and the second on a 12 core / 24 thread 1.8 GHz SoC. Both SoCs are available from Freescale employing 64-bit Power ISA. These SoCs meet our μ Server definition: integrate the entire server motherboard into a single microchip, except DRAM, NOR-boot and power conversion circuitry. Our focus is on density yet also on compute, memory and I/O BW balance. Our innovative hot-water based cooling infrastructure also supplies the electrical power to our compute node board, which enables us to put up to 128 compute nodes in a single 2U rack-unit. This unit will provide 3000 HW threads, 6 TB of DRAM and 32 Ethernet interfaces at 40Gbps. For SC14 we would demo single compute node air-cooled systems running Fedora20 and DB2.

Persistent In-Memory Computing with Emerging Non-Volatile Memory

Qingsong Wei, Jun Yang, Cheng Chen, Mingdi Xue, Chundong Wang, Renuga Kanagavelu (Data Storage Institute)

Emerging non-volatile Memory (NVM) is changing the way we think about memory and storage because they provide DRAM-like performance, disk-like high density and persistence. Current system is designed and optimized for traditional

DRAM-Disk hierarchy. When memory is replaced with NVM, memory is becoming persistent. The inefficiency of current architecture and system software significantly limits the potential of NVM. The objective of this project is to redesign the system and address architectural challenges and develop a suite of techniques which enable persistent in-memory computing with NVM for high performance computing. The major works we have achieved are NVM-accelerated File System and Consistent and Durable In-memory Key-value Store. NVM-accelerated File System explores NVM's byte-addressability and persistency to improve metadata latency and reduce metadata I/O. Consistent and Durable In-memory Key-value Store achieves high performance through designing NVM-based indexing data structure with low overhead of maintaining data consistency.

The AweSim Appkit: A Platform for Developing Web-Based HPC Applications

Alan Chalker, Dave Hudak (Ohio Supercomputer Center)

The Ohio Supercomputer Center's AweSim program is working to lower the barriers to entry for HPC based modeling and simulation. This initiative is designed to simplify the use of advanced simulation-driven design with integrated point-and-click apps to dramatically lower the cost of such simulations. AweSim uses four key technologies: (1) advanced models for manufacturing processes (2) a simple web-based user interface, (3) cloud-based computing resources and (4) computational solver software. The App Kit is a multilayer architecture to tie these components together, consisting of a web access layer, a web application layer and a system interface layer. The web access layer provides integrated single sign-on authentication with the App Store and industry-standard web proxy services. The web application layer is built on industry-standard technologies and contains reusable app templates. Finally, the system interface layer provides direct access to local resources such as databases, file systems and HPC job queues.

A New Node Design for HPC

Roger Ronald, Kevin Moran (System Fabric Works), Thomas Sohmers, Paul Sebexen (Rex Computing)

Current trends in processor technology have led to bottlenecks in memory latency, data movement, and storage. A need exists to accelerate practical exascale computing by focusing on power efficiency, simplicity of component design, and hierarchically designed systems. The only way to build a truly high performance computing system capable of meeting exascale demands and beyond is to design the simplest and thus most efficient processing unit, without sacrificing functionality or scalability.

System Fabric Works (SFW) and REX Computing have been jointly developing a full node architecture centered around

a new processor approach that is optimized for HPC. We propose to build a grid of many (potentially thousands of) Multiple Instruction, Multiple Data (MIMD) cores to execute instructions independently, drastically reducing the demands on memory bandwidth and thus greatly increasing the efficiency at which typical HPC workloads may be performed.

DEEP Collective Offload

Florentino Sainz, Vicenç Beltran, Jesús Labarta (Barcelona Supercomputing Center), Carsten Clauss, Thomas Moschny (ParTec Cluster Competence Center GmbH), Norbert Eicker (Juelich Research Center)

Exascale performance requires a level of energy efficiency only achievable with specialized hardware. Hence, the trend is clearly going towards heterogeneous architectures. However, in order to fully exploit the advantages of such heterogeneous hardware, there is a need for appropriate programming models that do not increase the complexity of the applications. In the DEEP project, we have extended the OmpSs programming model to offload MPI kernels dynamically by utilizing MPI's spawning feature, here namely provided by ParaStation MPI as the underlying infrastructure. That way, offloading to specialized hardware gets dramatically simplified while hiding the low-level and error prone MPI spawn calls. These facts together with its portability make this concept definitely interesting for a broader audience of the HPC community.

FPGA Development Tools and HMC/ FPGA Architectures

Steve Wallach (Convey Computer Corporation)

Convey has developed an OpenMP compiler that accepts standard OpenMP syntax and outputs a thread level intermediate language – HT (Hyper-Threading Tools). The HT tool sets, further compiles the intermediate language to Verilog.

The system developed can be used to create “FPGA personalities” for Convey developed hardware platforms. These platforms range from PCI cards, complete compute nodes, and systems that use HMC memory architectures

Blamange : Fat Binaries on a Diet

Kurt Keville (Massachusetts Institute of Technology), Anthony Skjellum (Auburn University), Mitch Williams (Sandia National Laboratories)

The latest generation of SoCs present new opportunities for performance and reliability by integrating multiple architectures in silica. At SC14 we intend to demo a cluster that takes advantage of a novel functionality within CUDA 6.0 that allows quick pivots and context switches within shared memory. This demo will show new strategies for hardware-level cyber-defense. The takeaway from this presentation is how the application is facilitated through fat binaries, with its broad implications to HPC.

Thursday, November 20

Emerging Technologies Exhibits

10am-3pm

Room: Show Floor - 233 Exhibits

HPC Impact Showcase Presentations

Chair: David Halstead (National Radio Astronomy Observatory)

10:20am-1pm

Room: Show Floor - 233 Theater

Is Healthcare Ready for HPC (and Is HPC Ready for Healthcare)?

Patricia Kovatch (Mount Sinai Hospital)

HPC is readily accepted (and perhaps expected) to tackle certain kinds of scientific questions in quantum chemistry, astrophysics and earthquake simulations. But how easy is it to apply typical HPC “know how” to healthcare and medical applications? What kinds of help can HPC provide and what are the barriers to success?

Biomedical researchers and clinicians are open to new tools and approaches to advance the diagnosis and treatment of human disease. One such approach is engaging interdisciplinary teams of computer scientists, applied mathematicians, computational scientists to work towards and improved understanding of the biological and chemical mechanisms behind disease.

Although progress is nascent, the opportunities and benefits of applying HPC resources and techniques to accelerate biomedical research and improve healthcare delivery are clear. One example of early success is the use of molecular dynamics to better understand protein interactions to design non-addictive painkillers. Clinicians are also using HPC applications for 3D visualizations of patient-specific surgical preparation for neurosurgery and cardiac surgery. State-of-the-art personalized and precision medicine is based on genomic sequencing and analysis on HPC resources.

High Performance Computing: An Essential Resource for Oil and Gas Exploration Production

Henri Calandra (Total)

For several decades the Oil and Gas industry has been continuously challenged to produce more hydrocarbons in response to the growing world demand for energy. Finding new economical oil traps has become increasingly challenging. The industry must invest in the design and definition of new advanced exploration and production tools. Application of HPC to oil and gas has dramatically increased the effectiveness of seismic exploration and reservoir management. Significantly

enhanced computational algorithms and more powerful computers have provided a much better understanding of the distribution and description of complex geological structures, opening new frontiers to unexplored geological areas. Seismic depth imaging and reservoir simulation are the two main domains that take advantage of the rapid evolution of high performance computing. Seismic depth imaging provides invaluable and highly accurate subsurface images reducing the risk of deep and ultra-deep offshore seismic exploration. The enhanced computational reservoir simulation models optimize and increase the predictive capabilities and recovery rates of the subsurface assets. With an order of magnitude increase in computing capability reaching an exaflop by the end of this decade, the reduction of computing time combined with rapid growth in data sets and problem sizes will provide far richer information to analyze. While next generation codes are developed to give access to new information of unrivaled quality, these codes will also present new challenges in building increasingly complex and integrated solutions and tools to achieve the exploration and production goals.

HPC - Accelerating Science through HPC

Dee Dickerson (Dow Chemical Company)

Today Dow has over 5,000 CPU cores used for a multitude of Research projects, Manufacturing Process Design, Problem Solving, and Optimization of existing Processes. High Performance Computing (HPC) has also enabled better materials design, models for plant troubleshooting especially in a situation where time is the main constraint.

HPC has enabled Dow researchers to meet stringent timelines and deliver implementable solutions to businesses. HPC provides the platform where large complex reacting flow models can be analyzed in parallel to significantly shorten the delivery time. The science and engineering community at Dow continues to advance technologies and develop cutting-edge computational capabilities involving some of the most complex multiphase reactive processes in dispersion applications. Such modeling capability development is only possible with the advancement of high speed, high capacity, and large memory HPC systems that are available to the industry today.

This presentation will show case the advancements Dow Researchers have made between using HPC and Lab Experiments.

How HPC Matters to Dresser Rand

Ravichandra Srinivasan (Dresser-Rand)

Abstract: TBA

Thursday, November 20

The Exascale Effect: the Benefits of Supercomputing Investment for U.S. Industry

Cynthia McIntyre, Chris Mustain (Council on Competitiveness)

The Council on Competitiveness has worked for over a decade to analyze, promote and strengthen America's use of advanced computing for economic competitiveness. The Council brings together top high performance computing leaders in industry, academia, government and the national laboratories. We know from experience that to out-compete is to out-compute.

Supported by a grant from the U.S. Department of Energy, the Council recently issued the report Solve – the Exascale Effect: the Benefits of Supercomputing Investment for U.S. Industry. The report takes a fresh look at: (1) the value of HPC to U.S. industry, (2) key actions that would enable companies to leverage advanced computing more effectively, and (3) how American industry benefits directly and indirectly from government investment at the leading edge of HPC.

The leaders of the Council's HPC Initiative, Senior Vice President Cynthia McIntyre and Vice President Chris Mustain, will share the findings of the report.

High Performance Computing: an Essential Resource for Oil & Gas Exploration Production

Henri Calandra (Total)

For several decades, the oil and gas industry has been continuously challenged to produce more hydrocarbons in response to the growing world demand for energy. Finding new economical oil traps has become increasingly challenging. The industry must invest in the design and definition of new advanced exploration and production tools. Application of HPC to oil and gas has dramatically increased the effectiveness of seismic exploration and reservoir management. Significantly enhanced computational algorithms and more powerful computers have provided a much better understanding of the distribution and description of complex geological structures, opening new frontiers to unexplored geological areas. Seismic depth imaging and reservoir simulation are the two main domains that take advantage of the rapid evolution of high performance computing. Seismic depth imaging provides invaluable and highly accurate subsurface images reducing the risk of deep and ultra-deep offshore seismic exploration. The enhanced computational reservoir simulation models optimize and increase the predictive capabilities and recovery rates of the subsurface assets. With an order of magnitude increase in computing capability reaching an exaflop by the end of this decade, the reduction of computing time combined with rapid growth in data sets and problem sizes will provide far richer information to analyze. While next generation codes are developed to give access to new information of unrivaled

quality, these codes will also present new challenges in building increasingly complex and integrated solutions and tools to achieve the exploration and production goals.

HPC of the Living Dead

Terry Boyd (Centers for Disease Control)

This session will use a fictitious “zombie virus” outbreak to demonstrate the multiple applications of HPC to public health. These include: (1) Big Data analysis to detect outbreaks; (2) spatial modeling and modeling of potential outbreaks to develop emergency response plans; (3) genomic evaluation of suspected disease outbreaks; (4) drug manufacturing and supply modeling; (5) contact tracing and response modeling to evaluate current activities and improve response activities; and (6) post-event analysis to prepare for the next epidemic.

Though this session will use a fictitious zombie outbreak as the basis for discussion in order to increase interest in the session, the uses described will be applicable to real-life outbreaks such as Ebola and H1N1.

Emerging Technologies Presentations

2pm-3pm

Room: Show Floor - 233 Theater

Tools and Techniques for Runtime Optimization of Large Scale Systems

Martin Schulz, Abhinav Bhatale, Peer-Timo Bremer, Todd Gamblin, Kathryn Mohror, Barry Rountree (Lawrence Livermore National Laboratory)

Future HPC systems will exhibit a new level of complexity, in terms of their underlying architectures and system software. At the same time, the complexity of applications will rise sharply, to both enable new science and exploit the new hardware features. Users will expect a new generation of tools and runtime support techniques that help address these challenges, work seamlessly with current and new programming models, scale to the full size of the machine, and address new challenges in resilience and power-aware computing.

We present a wide range of tool and runtime system research efforts targeting performance analysis, performance visualization, debugging and correctness tools, power-aware runtimes, and advanced checkpointing. These efforts are supported by activities on tool and runtime infrastructures that enable us to work with modular tool designs and support rapid tool prototyping. This “emerging technologies” booth will showcase our approaches to users, system designers and other tool developers.





HPC Interconnections

HPC Interconnections

HPC brings together people, ideas, and communities. And it is through the power of connections so evident at SC14 that some of HPC's greatest contributions to life today got their first exposure and traction. Building on the power of community, SC14 will focus on the increasing integration of HPC into modern life through HPC Interconnections – an enhanced version of the long-standing community elements of SC. HPC Interconnections will help everyone get more out of the conference, providing programs for everyone interested in building a stronger HPC community, including students, educators, researchers, international attendees and underrepresented groups.

HPC Interconnections comprises Broader Engagement, The Mentor/Protégé program, and programs for students, including the Student-Postdoc Job & Opportunity Fair, the Student Cluster Competition, Experiencing HPC for Undergraduates and Student Volunteers. The goal is to ensure that all attendees make meaningful connections to help them during SC14 and after. The HPC Interconnections program appreciates the support of the following sponsors:

- Bank of America
- Chevron
- DataDirect Networks
- The Walt Disney Co.
- Geist
- General Motors
- Lawrence Livermore National Laboratory
- MathWorks
- National Security Agency
- Procter & Gamble
- Schlumberger

HPC Interconnections

HPC Interconnections comprises Broader Engagement and programs for Students including the Student-Postdoc Job & Opportunity Fair, Student Cluster Competition, Experiencing HPC for Undergraduates, and Student Volunteers

Saturday, November 15

Broader Engagement (BE) Orientation

Chair: Mary Ann Leung (Sustainable Horizons Institute)

4pm-5pm

Room: 288-89

SC14 Broader Engagement Welcome

Mary Ann Leung (Sustainable Horizons Institute),
Elizabeth Leake (University Wisconsin-Milwaukee),
Raquell Holmes (improvscience)

This session will begin the SC14 Broader Engagement program with a welcome from the Chair and Deputy Chair. An overview of the program activities will be presented and important topics guiding the week's activities will be discussed. The session will also include an interactive session led by Raquell Holmes from improvscience and focused on improving your ability to collaborate, communicate and work with each other to fully take advantage of the conference and each other's expertise.

BE Submitting Posters: Making Your Research Known

Chair: Richard Coffey (Argonne National Laboratory)

5pm-6:30pm

Room: 288-89

BE Panel and Discussion: Submitting Posters: Making Your Research Known

Richard Coffey (Argonne National Laboratory), Michela Taufer (University of Delaware), Raquell Holmes (improvscience), Tony Drummond (Lawrence Berkeley National Laboratory)

This session is designed to help BE participants learn why and how to submit and present posters in the main technical track of the SC15 conference. The session will motivate BE participants by demystifying and demonstrating the importance to their careers of presenting a poster at a major conference. The goal is to inspire BE participants to start the journey to submissions now for SC15. A panel of past BE participants and small group exercises will cover the following topics:

- Situate the SC posters in the SC publishing ecosystem
- Tips and tricks for writing and developing the poster
- Lead participants through the content submission process

- Guidance on both the peer review process and poster competition
- Stories of successful BE submissions
- Share information about opportunities that remove roadblocks: funding, support, etc.

Sunday, November 16

BE Plenary I: Introduction to HPC and Its Applications

Chair: Tony Drummond (Lawrence Berkeley National Laboratory)

8:30am-10am

Room: 288-89

An Introduction to High Performance Computing and the SC Conference

William Gropp (University of Illinois at Urbana-Champaign)

What is HPC and what happens at the SC conference? This talk will discuss the different kinds of HPC, what HPC is used for, and how people work in the field of HPC. The premier conference in HPC is SC, and I will also talk about the history of the conference, the role it plays in the field of HPC, and provide a behind-the-scenes peek at how SC comes together every year.

Introduction to the On-Line Course "CS267: Applications of Parallel Computers"

James Demmel

CS267 is taught every spring to graduate students from many departments at UC Berkeley, and since 2012 it has also been offered on-line by XSEDE to students around the world, with free NSF supercomputer accounts to do the autograded homework. After outlining the course, we go into more detail on two overarching themes: (1) How to recognize common computational patterns that arise nearly all high performance applications, and use them to make implementations both more efficient and productive. (2) Since communication, i.e. moving data, either between levels of a memory hierarchy or processors over a network, is much more expensive than arithmetic, we present algorithms for a number of these patterns that minimize communication. Finally, we discuss other on-line course offerings.

BE Session IA: HPC Memory Lane and Future Roadmap (Advanced)

Chair: Tony Drummond (Lawrence Berkeley National Laboratory)

10:30am-11:15am

Room: 288-89

Overview of the Exascale Challenges

John Shalf (Lawrence Berkeley National Laboratory)

The talk will provide an overview of the challenges posed by physical limitations of the underlying silicon-based CMOS technology, introduce the next generation of emerging machine architectures and the anticipated effect on the way we program machines in the future.

For the past 25 years, a single model of parallel programming (largely bulk-synchronous MPI), has for the most part been sufficient to permit translation of this into reasonable parallel programs for more complex applications. In 2004, however, a confluence of events changed forever the architectural landscape that underpinned our current assumptions about what to optimize for when we design new algorithms and applications. This talk will describe the challenges of programming future computing systems and provide some highlights from the search for durable programming abstractions more closely track emerging computer technology trends so when we convert our codes over, they will last through the next decade.

BE Session IB: What is HPC and Why is it Relevant? What is SC? (Intro)

Chair: Karla Morris (Sandia National Laboratories)

10:30am-11:15am

Room: 290

HPC as a State of Mind

James Costa (Sandia National Laboratories)

The term High Performance Computing (HPC) is used everywhere, has a certain cache to it and is difficult to rigorously define solely in terms of technical parameters. We can think of HPC in terms of operations speed, architectures, lines of code, power consumption or a named computer with a position on the Top 500 List. We link the term to the now almost forgotten concept of a supercomputer. Last, there is the unavoidable fact that today's accepted HPC defining metric will soon become tomorrow's definition of commodity computing, or worse, your new laptop. This talk will probe a little deeper into this concept of HPC while showing how SC14 can provide you a path for personal exploration. And for those who really want to go deep, we might even examine the humor of a conference dedicated to HPC and the 'supercomputer' but has neither item in its official name. #hpcmatters

BE Session IA: Multicore Programming (Advanced)

Chair: Tony Drummond (Lawrence Berkeley National Laboratory)

11:15am-12pm

Room: 288-89

Designing Dense Linear Algebra Libraries for High Performance Computing

Jack Dongarra (University of Tennessee, Knoxville)

In this talk we will look at some of the software and algorithm challenges we face when designing numerical libraries at scale. These challenges, such as management of communication and memory hierarchies through a combination of compile-time and run-time techniques, as well as the increased scale of computation, depth of memory hierarchies, range of latencies, and increased run-time environment variability, require a different approach to what has been done in the past.

BE Session IB: HPC Memory Lane and Future Roadmap (Intro)

Chair: Karla Morris (Sandia National Laboratories)

11:15am-12pm

Room: 290

Overview of the Exascale Challenges

This talk will provide an overview of the challenges posed by the physical limitations of the underlying silicon based CMOS technology, introduce the next generation of emerging machine architectures, and the anticipated effect on the way we program machines in the future.

For the past twenty-five years, a single model of parallel programming (largely bulk-synchronous MPI), has for the most part been sufficient to permit translation of this into reasonable parallel programs for more complex applications. In 2004, however, a confluence of events changed forever the architectural landscape that underpinned our current assumptions about what to optimize for when we design new algorithms and applications. We have been taught to prioritize and conserve things that were valuable 20 years ago, but the new technology trends have inverted the value of our former optimization targets. The time has come to examine the end result of our extrapolated design trends and use them as a guide to re-prioritize what resources to conserve in order to derive performance for future applications. This talk will describe the challenges of programming future computing systems. It will then provide some highlights from the search for durable programming abstractions more closely track emerging computer technology trends so that when we convert our codes over, they will last through the next decade.

BE Session II: HPC Applications and Q&A*Chair: Tony Drummond (Lawrence Berkeley National Laboratory)***1:30pm-3pm****Room: 288-89****Enabling Physics Insight with Multidimensional Computer Simulations***Alice Koniges (Lawrence Berkeley National Laboratory)*

Physics and numerical simulations have progressed hand in hand since the beginning of computer science. With the ever-increasing compute power offered at the highest end, numerical simulation of physics experiments yields both design and understanding at an unprecedented level. I will describe some recent physics simulations in the field of accelerator and fusion energy research and discuss how they rely on HPC for true fidelity, design and predictability. I will also discuss how the latest trends in HPC are affecting simulation techniques.

QRing – A Scalable Parallel Software Tool for Quantum Transport Simulations in Carbon Nanoring Devices Based on NEGF Formalism and a Parallel C++ / MPI / PETSc Algorithm*Mark A. Jack (Florida A&M University)*

The ability of a nanomaterial to conduct charge is essential for many nanodevice applications. Numerous studies have shown that disorder, including defects and phononic or plasmonic effects, can disrupt or block electric current in nanomaterials. An accurate theoretical account of both electron-phonon and electron-plasmon coupling in carbon nanotube- and nanoring-based structures is key in order to properly predict the performance of these new nanodevices. The parallel sparse matrix library PETSc provides tools to optimally divide memory and computational tasks required in inverting the Hamiltonian matrix across multiple compute cores or nodes. The results at different energy levels are integrated via a second layer of parallelism to obtain integrated observables. Because the inversion at each energy step dominates the run time and memory use of the code, it is important that the code scales well to realistic system sizes. Several different direct and iterative linear solvers and libraries were compared for their scalability and efficiency, including the Intel MKL shared-memory, direct dense solver, the MULTifrontal Massively Parallel direct Solver (MUMPS), and iterative sparse solvers bundled with PETSc. The NSF XSEDE resource ‘Stampede’ at the Texas Advanced Computing Center as well as regional resources in Florida (SSERCA) are used for development and benchmarking, while DOE OLCF computational resources ‘Titan’ and ‘NICS/Beacon’ are used for physics production runs.

Big Data, Page Ranking and Epidemic Spread Modeling*Nahid Emad (Versailles Saint-Quentin-en-Yvelines University)*

The surge of medical and nutritional data in the field of health requires research of models and methods as well as the development of data analysis tools. The spread of infectious diseases, detection of biomarkers for prognosis and diagnosis of the disease, research of indicators for personalized nutrition and/or medical treatment are some typical examples of problems to solve. In this talk, we focus on the spread of contagious diseases and show how the eigenvalue equation intervenes in models of infectious disease propagation and could be used as an ally of vaccination campaigns in the actions carried out by health care organizations. The stochastic model based on PageRank allows simulating the epidemic spread, where a PageRank-like infection vector is calculated to help establish efficient vaccination strategy. Due to the size and the particular structure of underlying social networks, this calculation requires considerable computational resources as well as storage means of very large quantity of data and represents a big challenge in high performance computing. The computation methods of PageRank in this context are explored. The experiments take into account very large network of individuals imposing the challenging issue of handling very big graph with complex structure. The computational challenges of some other applications such as identification of biomarkers for prognosis and diagnosis of a disease will also be discussed.

Extreme-Scale In-Situ Data Analysis*Janine C. Bennett (Lawrence Berkeley National Laboratory)*

Steady improvements in computing resources enable ever more enhanced simulations, but data input/output (I/O) constraints are impeding their impact. Despite increases in temporal resolution, the gap between time steps saved to disk keeps increasing as computational power continues to outpace I/O capabilities. Consequently, we are seeing a paradigm shift away from the use of prescribed I/O frequencies and post-process-centric data analysis toward a more flexible concurrent paradigm in which raw simulation data is processed as it is computed. This talk will survey applied mathematics and computer science research challenges introduced by this paradigm shift and highlight active research efforts to deal with these issues.

Q&A

Tony Drummond (Lawrence Berkeley National Laboratory)

BE Session III: Careers & Diversity

Chair: Mary Ann Leung (Sustainable Horizons Institute)

3:30pm-4:30pm

Room: 288-89

Careers and Diversity

Dot Harris (Department of Energy)

BE Session III: Introduction to Networking and Posters

Chair: Tony Drummond (Lawrence Berkeley National Laboratory)

4:30pm-5pm

Room: 288-89

Introduction to the Networking Event and Posters

*Tony Drummond (Lawrence Berkeley National Laboratory),
Karla Morris (Sandia National Laboratories)*

An introduction to the networking event and poster session.

BE Networking Event: Student/Faculty/Professional Posters

Chair: Tony Drummond (Lawrence Berkeley National Laboratory)

5:30pm-8pm

Room: 288-89-90

BE participants are invited to present a technical poster. The posters will be reviewed by seasoned professionals who will also provide guidance, suggestions and mentoring for improving posters and poster presentations. Please note that this event is open only to invited participants.

Monday, November 17**BE Plenary II: Big Data and Exascale Challenges**

Chair: Karla Morris (Sandia National Laboratories)

8:30am-10am

Room: 288-89

Big Data and Combustion Simulation

Jacqueline H. Chen (Sandia National Laboratories)

Exascale computing will enable combustion simulations in parameter regimes relevant to next-generation combustors burning alternative fuels needed to provide the underlying science base for developing predictive combustion models and the analysis of massive volumes of data. It will enable combustion research with complex fuels in high-pressure, turbulent

environments associated with IC engines and gas turbines for power generation. The current practice at the petascale of writing out raw data from simulations to persistent disk and re-reading in the data later for analysis won't scale at the exascale due to the disparity among computational power, machine memory and I/O bandwidth. It is anticipated that much of the current analysis and visualization will need to be integrated with the combustion simulation in situ in a comprehensive overall work flow on an exascale machine. In addition to understanding the specific performance characteristics of individual solver, analysis and visualization algorithms, we need to understand the overall workload and impact of various designs in the context of the end-to-end workflow. Finally, programming models and runtimes that optimize performance for complex in situ workflows will be discussed.

Lessons from Protein Folding

Silvia Crivelli (Lawrence Berkeley National Laboratory)

Protein folding is one of the biggest challenges in modern biology with both theoretical and practical significance. It is about finding the shape of proteins, which is key to understanding the mechanisms of life, to finding new drugs to combat disease, and to designing new proteins with desired functions not currently found in nature. Diseases associated with proteins not working properly include cancer, Alzheimer's, and mad cow disease. The complexity of the protein-folding problem requires a multidisciplinary and community wide approach. To this end, scientists have tried different social-based approaches to advancing protein folding: They have tried a competition called CASP (Critical Assessment of techniques for protein Structure Prediction) and they have developed a computer game to harness the human intuition to guess the correct shapes. More recently, we have begun using social media and science gateways as a way to bring together labs and individuals from all over the world to a "virtual discussion table". These large-scale efforts have created new challenges that require HPC and data science. In this talk I'll discuss some lessons that we've learned and some of the challenges that still remain.

BE Session IVA: Programming for Exascale - Advanced Sessions

Chair: Karla Morris (Sandia National Laboratories)

10:30am-12pm

Room: 288-89

Profiling HPC Codes with TAU

Sameer Shende (University of Oregon)

The TAU Performance System helps computational scientists evaluate the performance of their codes. This talk will highlight the features of the toolkit and show how its profiling and debugging capabilities can be applied to isolate performance and memory bugs. The talk will introduce basic profiling and

performance tuning concepts and will show new techniques that include event-based sampling, rewriting binaries, runtime preloading, OpenMP instrumentation using OMPT tools interface, and source based instrumentation.

TAU supports both profiling and tracing modes of measurement. This talk will highlight TAU's use for a wide variety of languages and runtime systems, including MPI, CUDA, pthread, OpenMP and OpenACC.

Common Computational Kernels and Motifs

Erich Strohmaier (Lawrence Berkeley National Laboratory)

Abstract: TBD

Programming Models in HPC Systems and Application Development and Tuning Strategies

Khaled Ibrahim (Lawrence Berkeley National Laboratory)

The architectural trends in HPC systems show increase in heterogeneity, decrease in physical memory per core, increase in core count per node, and the possibility of having slower cores. These trends exacerbate the need to support hybrid programming models within the same system. In this talk, we visit the dominant programming models in HPC systems including the MPI two-sided, the PGAS one-sided, and the accelerator-based models. We discuss how these programming models influence the application development strategy, for instance for maintaining data consistency, structuring communication, decomposing tasks and managing the data layout.

BE Session IVB: Parallel Programming - Introduction Sessions

Chair: *Tony Drummond (Lawrence Berkeley National Laboratory)*

10:30am-12pm

Room: 290

HPC Hardware Overview and the TOP500

Erich Strohmaier (Lawrence Berkeley National Laboratory)

Programming Models and Applications in HPC Systems
Khaled Ibrahim (Lawrence Berkeley National Laboratory)
The architectural trends in HPC systems show increase in heterogeneity, decrease in physical memory per core, increase in core count per node, and the possibility of having slower cores. These trends exacerbate the need to support hybrid programming models within the same system. In this talk, we visit the dominant programming models in HPC systems including the MPI two-sided, the PGAS one-sided, and the accelerator-based models. We discuss how these programming models influence the application development strategy, for instance for maintaining data consistency, structuring communication, decomposing tasks, and managing the data layout.

Programming Models and Applications in HPC Systems

Khaled Ibrahim

The architectural trends in HPC systems show increase in heterogeneity, decrease in physical memory per core, increase in core count per node, and the possibility of having slower cores. These trends exacerbate the need to support hybrid programming models within the same system. In this talk, we visit the dominant programming models in HPC systems including the MPI two-sided, the PGAS one-sided, and the accelerator-based models. We discuss how these programming models influence the application development strategy, for instance for maintaining data consistency, structuring communication, decomposing tasks, and managing the data layout.

An introduction to Profiling HPC Codes with TAU

Sameer Shende (University of Oregon)

The TAU Performance System helps computational scientists evaluate the performance of their codes. This talk will highlight the features of the toolkit and show how its profiling and debugging capabilities can be applied to isolate performance and memory bugs. The talk will introduce basic profiling and performance tuning concepts and will show new techniques that include event-based sampling, rewriting binaries, runtime preloading, OpenMP instrumentation using OMPT tools interface, and source based instrumentation.

TAU supports both profiling and tracing modes of measurement. The talk will highlight TAU's use for a wide variety of languages and runtime systems, including MPI, CUDA, pthread, OpenMP and OpenACC.

BE Mentor/Protégé Session & Mixer

Chair: *Christine Sweeney (Los Alamos National Laboratory)*

1:30pm-3:30pm

Room: 288-89

Mentor-Protégé Session and Mixer

Christine Sweeney (Los Alamos National Laboratory)

The Mentor-Protégé event supports participants in the HPC Interconnections program by creating opportunities for attendees to establish developmental professional relationships. By pairing fairly new SC participants with more experienced attendees, the Mentor-Protégé program provides introduction, engagement and support. The Mentor-Protégé Mixer begins with a brief talk by a mentor and protege, then Mentor-Protégé pairs move on to discuss both networking and technical development pathways to help newcomers to the conference integrate with the larger SC community. The event will also include some time at the end for proteges to network with other mentors in attendance. Refreshments will be served. Questions on the program or event can be directed to mentor-protége@info.supercomputing.org.

Experiencing HPC For Undergraduates Orientation

Chair: Jeffrey K. Hollingsworth (University of Maryland)

3pm-5pm

Room: 290

Experiencing HPC For Undergraduates Orientation

Jeffrey K. Hollingsworth (University of Maryland), Alan Sussman (University of Maryland)

This session provides an introduction to HPC for the participants in the HPC for Undergraduates Program. Topics will include an introduction to MPI, shared memory programming, domain decomposition and the typical structure of scientific programming.

BE Programming Challenge

Chair: Kenneth Craft (Intel Corporation)

3:30pm-7pm

Room: 288-89

Programming Challenge

Kenneth Craft (Intel Corporation)

The Broader Engagement programming challenge is an opportunity for new programmers and seasoned programmers within the BE program to compete against each other in friendly competition. With multiple levels and multiple languages, there is something for everyone to excel at. You only need your laptop to compete, as a virtual machine will be provided with all the needed software needed. Participants can also use their own software if they prefer. Test your abilities by writing legible code in a timely manner. Top competitors will receive prizes.

improvscience Programming Challenge

Raquell Holmes (improvscience), Ritu Arora (University of Texas at Austin)

Real world solutions are often produced by teams whose members' have diverse skills being combined to address a unique problem. The "improvscience Programming Challenge" is a problem-solving and a team-building activity that is being organized for the first time as a part of the BE program. Participants will work in small teams to developing a solution strategy for an HPC or a big data management problem. The submissions will be judged on the basis of the technical merit, creativity of the approach and engagement of the team members in developing the solution strategy. The challenge is organized in three stages: team formation, initial solution and revisions, and final presentation.

First Time at SC? Here's How to Get Oriented!

Chair: Bruce Loftis (University of Colorado - Boulder)

6pm-6:30pm

Room: 286-87

First Time at SC? Here's How to Get Oriented

Bruce Loftis (University of Colorado Boulder)

This session is for those who are attending the SC conference for the first time. Presenters will briefly describe the various parts of the conference, make some suggestions about what events you should attend and answer your questions. If you are a returning attendee and know others who may benefit from such a session, please help spread the word.

Student Cluster Competition Kickoff

Chair: Dustin Leverman (Oak Ridge National Laboratory)

7:30pm-9pm

Show Floor - 410

Student Cluster Competition Kickoff

Daniel Kamalic (GNS Healthcare)

Come help kick off the Student Cluster Competition and cheer on the 12 student teams competing in this real-time, non-stop, 48-hour challenge to assemble a computational cluster on the SC14 exhibit floor and demonstrate the greatest sustained performance across a series of applications. Teams selected for this year's competition come from universities all over the United States, Germany, Taiwan, China, Singapore, and

Australia. The teams are using big iron hardware with off-the-shelf (or off-the-wall!) solutions for staying within a 26 amp power limit.

This year's competition is generously sponsored by Bank Of America, Chevron, Data Direct Networks (DDN), Geist Global, General Motors, MathWorks, Procter & Gamble and Schlumberger.

Tuesday, November 18

Student Cluster Competition

Chair: Daniel Kamalic (GNS Healthcare)

10am-6pm

Room: Show Floor - 410

Student Cluster Competition

Daniel Kamalic (GNS Healthcare)

Come to Booth 410 on the show floor to meet the 12 student teams competing in this real-time, non-stop, 48-hour challenge to assemble a computational cluster on the SC14 exhibit floor and demonstrate the greatest sustained performance across a series of applications. Teams selected for this year's competition come from universities all over the United States, Germany, Taiwan, China, Singapore, and Australia. The teams are using big iron hardware with off-the-shelf (or off-the-wall!) solutions for staying within a 26 amp power limit.

This year's competition is generously sponsored by Bank Of America, Chevron, Data Direct Networks (DDN), Geist Global, General Motors, MathWorks, Procter & Gamble and Schlumberger.

Experiencing HPC For Undergraduates: Introduction to HPC Research

Chair: Jeffrey K. Hollingsworth (University of Maryland)

10:30am-12pm

Room: 295

Experiencing HPC For Undergraduates: Introduction to HPC Research

*Jacqueline Chen (Sandia National Laboratories),
William Kramer (University of Illinois at Urbana-Champaign),
John Mellor-Crume (Rice University), Minish Parashar
(Rutgers University)*

A panel of leading practitioners in HPC will introduce the various aspects of HPC, including architecture, applications, programming models and tools.

BE Session: Interviewing & Mock Interviews

Chair: Christine E. Harvey (MITRE Corporation)

1:30pm-3pm

Room: 288-89

Interviewing Panel and Mock Interviews

Christine Harvey (MITRE Corporation), Jose Castillo (San Diego State University), Heidi Lee Alvarez (Florida International University), Elizabeth Bautista (Lawrence Berkeley National Laboratory), Terry Boyd (Centers for Disease Control), Richard Coffey (Argonne National Laboratory)

Interviewing is a vital part of any career and is often a new and challenging experience for students and those looking to enter the workforce. This session will be split into two parts, a panel on interviewing followed by the opportunity to partake in mock interviews. The panel will contain members from a variety of organizations including industry, government, and academia and discuss important topics in interviewing for those new to the practice. Experts on the panel will discuss best practices in interviewing especially in the highly technical area of HPC. The opening part of the session will cover general interviewing and then allow time for questions. The second half of the session will match participants with volunteer interviewers for a short series of mock interviews. These short interviews will give attendees a chance to gain one-on-one experience interviewing and allow participants to receive constructive feedback. Participants, bring your resumes!

BE Session: Guided Tours of the Poster Session

5pm-7pm

Room: New Orleans Theater Lobby

BE participants will take a guided tour of selected Research and ACM Student Research Competition Posters.

Wednesday, November 19

Student-Postdoc Job & Opportunity Fair

Chair: Beth McKown (National Center for Supercomputing Applications)

10am-3pm

Room: 288-89

Student-Postdoc Job & Opportunity Fair

This face-to-face event will be open to all students and postdocs participating in the SC14 conference, giving them an opportunity to meet with potential employers. There will be employers from research labs, academic institutions, recruiting agencies and private industry who will meet with students and postdocs to discuss graduate fellowships, internships, summer jobs, co-op programs, graduate school assistant positions and/or permanent employment.

Student Cluster Competition

Chair: Dustin Leverman (Oak Ridge National Laboratory)

10am-5pm

Room: Show Floor - 410

Student Cluster Competition

Daniel Kamalic (GNS Healthcare)

Come to booth 410 on the show floor to meet the 12 student teams competing in this real-time, non-stop, 48-hour challenge to assemble a computational cluster on the SC14 exhibit floor and demonstrate the greatest sustained performance across a series of applications. Teams selected for this year's competition come from universities all over the United States, Germany, Taiwan, China, Singapore, and Australia. The teams are using big iron hardware with off-the-shelf (or off-the-wall!) solutions for staying within a 26 amp power limit. This year's competition is generously sponsored by Bank Of America, Chevron, Data Direct Networks (DDN), Geist Global, General Motors, MathWorks, Procter & Gamble and Schlumberger.

Experiencing HPC For Undergraduates:

Graduate Student Perspective

Chair: Jeffrey K. Hollingsworth (University of Maryland)

10:30am-12pm

Room: 295

Experiencing HPC For Undergraduates:

Graduate Student Perspective

Kai Ren (Carnegie Mellon University), Ismail El-Helw (VU University Amsterdam), Adam T. McLaughlin (Georgia Institute of Technology), Maciej Besta (ETH Zurich), Dhairya Malhotra (University of Texas at Austin), Mehmet Can Kurt (Ohio State University)

This session will be held as a panel discussion. Current graduate students, all of whom are candidates for the Best Student Paper Award in the Technical Papers program at SC14, will discuss their experiences in being a graduate student in an HPC discipline. They will also talk about the process of writing their award-nominated paper.

Student Cluster Competition Grand Finale

Chair: Dustin Leverman (Oak Ridge National Laboratory)

5pm-6pm

Show Floor - 410

Student Cluster Competition Grand Finale

Daniel Kamalic (GNS Healthcare)

Come to booth 410 on the show floor to cheer on the 12 student teams as they turn in their final results in this real-time, non-stop, 48-hour challenge to assemble a computational cluster on the SC14 exhibit floor and demonstrate the greatest sustained performance across a series of applications. Teams selected for this year's competition come from universities all over the United States, Germany, Taiwan, China, Singapore, and Australia. The teams are using big iron hardware with off-the-shelf (or off-the-wall!) solutions for staying within a 26 amp power limit.

This year's competition is generously sponsored by: Bank Of America, Chevron, Data Direct Networks (DDN), Geist Global, General Motors, MathWorks, Procter & Gamble and Schlumberger.

Thursday, November 20

Experiencing HPC for Undergraduates

Chair: Jeffrey K. Hollingsworth (University of Maryland)

10:am-12pm

Room: 295

Experiencing HPC For Undergraduates:

Careers in HPC

David Bernholdt (Oak Ridge National Laboratory), Brad Chamberlain (Cray Inc.), Rob Schreiber (Hewlett Packard), Michela Taufer (University of Delaware)

This panel features a number of distinguished members of the HPC research community discussing their varied career paths. The panel includes representatives from industry, government labs and universities. The session will include ample time for questions from the audience.

HPCI Scavenger Hunt Awards

Chair: Mary Ann Leung (*Sustainable Horizons Institute*)

12pm-12:45pm

Room: 274

HPCI Scavenger Hunt Awards Session

Mary Ann Leung (Sustainable Horizons Institute)

This session will review the activities and results of the Scavenger Hunt and announce the winners.

Get involved in the Student Cluster Competition

Chair: Dustin Leverman (*Oak Ridge National Laboratory*)

1:30pm-3pm

Room: 295

Get involved in the Student Cluster Competition

Daniel Kamalic (GNS Healthcare)

Come share questions and feedback with teams and organizers from the Student Cluster Competition. Learn how your university, high school or company can be involved next year in the “Premier Event in Computer Sports.”

This year’s competition is generously sponsored by Bank Of America, Chevron, Data Direct Networks (DDN), Geist Global, General Motors, MathWorks, Procter & Gamble and Schlumberger.

BE Wrap Up

Chair: Mary Ann Leung (*Sustainable Horizons Institute*)

3:30pm-5pm

Room: 295

The Broader Engagement Program will conclude with this wrap-up session. We will review the week’s activities and solicit feedback from participants. Bring your ideas about what worked and did not work and how we can further the goals of the program at future SC conferences.

Keynote/ Invited Talks

Keynote/Invited Talks

At SC14 talks featured in years past such as Masterworks, Plenary talks, and State-of-the-Field, were combined under a single banner, called “Invited Talks”. We will continue this practice with renewed efforts to develop Invited Talks as a premier component of the Program for SC14.

Twelve Invited Talks will feature leaders in the areas of high-performance computing, networking and storage. Invited Talks will typically concern innovative technical contributions and their applications to address the critical challenges of the day.

Additionally, these talks will often concern the development of a major area through a series of important contributions and provide insights within a broader context and from a longer-term perspective. At Invited Talks, you should expect to hear about pioneering technical achievements, the latest innovations in supercomputing and data analytics, and broad efforts to answer some of the most complex questions of our time.

Keynote/Invited Talks

Tuesday, November 18

Keynote

Chair: Ron Benza (Benza Group, Inc.)

8:30am-10am

Room: New Orleans Theater



Dr. Brian Greene, physicist, string theorist, and best-selling author, will discuss what happens at the intersection of science, computing, and society.

Described by The Washington Post as “the single best explainer of abstruse concepts in the world today,” Brian Greene is one

of the world’s leading theoretical physicists and a brilliant, entertaining communicator of cutting-edge scientific concepts. Greene’s national bestseller, *The Elegant Universe*, recounts the theories of general relativity and quantum mechanics. The book sold over a million copies, and became an Emmy and Peabody Award-winning NOVA special. Greene’s second book, *The Fabric of the Cosmos*, spent six months on The New York Times Best Sellers list and was adapted into a NOVA miniseries on PBS. Greene’s latest bestseller, *The Hidden Reality: Parallel Universes and the Deep Laws of the Cosmos*, was published in January 2011.

Invited Talks

Chair: Robert F. Lucas (Information Sciences Institute)

10:30am-12pm

Room: New Orleans Theater

Quantum Computing Paradigms for Probabilistic Inference and Optimization

Masoud Mohseni (Google)

Bio: Masoud Mohseni is a Senior Research Scientist at Google Quantum Artificial Intelligence Laboratory, where he develops novel machine learning algorithms that fundamentally rely on quantum dynamics. A former research scientist and a principal investigator at the Research Laboratory of Electronics at MIT, Dr. Mohseni was also a scientific consultant at the Future and Emerging Technologies program of the European Commission on the LANDAUER Nanoscale Computing project and a consultant at BBN Technologies in the Disruptive Information Processing Group. He obtained his Ph.D in physics in 2007 from University of Toronto. He then moved to Harvard University to complete a postdoctoral program in quantum simulation. Dr. Mohseni’s current research addresses some of the fundamental problems at the interface of artificial intelligence, quantum computing, and physics of complex quantum systems. He has developed quantum algorithms for probabilistic inference

and optimization, scalable characterization and simulation of many-body systems, optimal and robust quantum transport in disordered open quantum systems, and quantum-assisted imaging. His algorithms have been successfully implemented in several leading quantum computing architectures. He has also made pioneering contributions to the theory of energy transfer in biomolecular networks. He is an author and the leading editor of the first scientific book on Quantum Effects in Biology that was recently published by Cambridge University Press.

Abstract: Over the past 30 years, several computational paradigms have been developed based on the premise that the laws of quantum mechanics could provide radically new and more powerful methods of information processing. One of these approaches is to encode the solution of a computational problem into the ground state of a programmable many-body quantum Hamiltonian system. Although, there is empirical evidence for quantum enhancement in certain problem instances, there is not a full theoretical understanding of the conditions for quantum speed up for problems of practical interest, especially hard combinatorial optimization and inference tasks in machine learning. In this talk, I will provide an overview of quantum computing paradigms and discuss the progress at the Google Quantum Artificial Intelligence Lab towards developing the general theory and overcoming practical limitations. Furthermore, I will discuss two algorithms that we have recently developed known as Quantum Principal Component Analysis and Quantum Boltzmann Machine.

Super Debugging: It Was Working Until I Changed ...

David Abramson (University of Queensland)

Bio: Professor David Abramson has been involved in computer architecture and high performance computing research since 1979. He has held appointments at Griffith University, CSIRO, RMIT and Monash University. Most recently at Monash he was the Director of the Monash e-Education Centre, Deputy Director of the Monash e-Research Centre and a Professor of Computer Science in the Faculty of Information Technology. He held an Australian Research Council Professorial Fellowship from 2007 to 2011. He has worked on a variety of HPC middleware components including the Nimrod family of tools and the Guard relative debugger. Abramson is currently the Director of the Research Computing Centre at the University of Queensland. He is a fellow of the Association for Computing Machinery (ACM) and the Academy of Science and Technological Engineering (ATSE), the Australian Computer Society and a Senior Member of the IEEE.

Abstract: Debugging software is a complex, error prone, and often a frustrating experience. Typically, there is only limited tool support, and many of these do little more than allow a user to control the execution of a program and examine its

run-time state. This process becomes even more difficult in supercomputers that exploit massive parallelism because the state is distributed across processors, and additional failure modes (such as race and timing errors) can occur. We have developed a debugging strategy called “Relative Debugging,” that allows a user to compare the run-time state between executing programs, one being a working, “reference” code, and the other being a test version. In this talk, I will outline the basic ideas in relative debugging, and will give examples of how it can be applied to debugging supercomputing applications.

Invited Talks

Chair: Rick Stevens (*Argonne National Laboratory, University of Chicago*)

3:30pm-5pm

Room: New Orleans Theater

Usable Exascale and Beyond Moore’s Law

Horst Simon (Lawrence Berkeley National Laboratory)

Bio: Horst Simon is an internationally recognized expert in computer science and applied mathematics and the Deputy Director of Lawrence Berkeley National Laboratory (Berkeley Lab). Simon joined Berkeley Lab in early 1996 as director of the newly formed National Energy Research Scientific Computing Center (NERSC) and was one of the key architects in establishing NERSC at its new location in Berkeley. Under his leadership NERSC enabled important discoveries for research in fields ranging from global climate modeling to astrophysics. Simon was also the founding director of Berkeley Lab’s Computational Research Division, which conducts applied research and development in computer science, computational science, and applied mathematics. In his prior role as Associate Lab Director for Computing Sciences, Simon helped to establish Berkeley Lab as a world leader in providing supercomputing resources to support research across a wide spectrum of scientific disciplines. He is also an adjunct professor in the College of Engineering at the University of California, Berkeley. In that role he worked to bring the Lab and the campus closer together, developing a designated graduate emphasis in computational science and engineering. In addition, he has worked with project managers from the Department of Energy, the National Institutes of Health, the Department of Defense and other agencies, helping researchers define their project requirements and solve technical challenges. Simon’s research interests are in the development of sparse matrix algorithms, algorithms for large-scale eigenvalue problems, and domain decomposition algorithms for unstructured domains for parallel processing. His algorithm research efforts were honored with the 1988 and the 2009 Gordon Bell Prize for parallel processing research. He was also member of the NASA team that developed the NAS Parallel Benchmarks, a widely used standard for evaluating the performance of massively parallel systems. He is co-editor of the biannual TOP500 list that tracks the most powerful supercomputers worldwide, as

well as related architecture and technology trends. He holds an undergraduate degree in mathematics from the Technische Universität, in Berlin, Germany, and a Ph.D. in Mathematics from the University of California at Berkeley.

Abstract: As documented by the TOP500, high performance computing (HPC) has been the beneficiary of uninterrupted growth, and performance of the top HPC systems doubled about every year until 2004, when Dennard scaling tapered off. This was based on the contributions of Moore’s law and the increasing parallelism in the highest end system. Continued HPC system performance increases were then obtained by doubling parallelism. However, over the last five years HPC performance growth has been slowing measurably, and in this presentation several reasons for this slowdown will be analyzed. To reach usable exascale performance over the next decade, some fundamental changes will have to occur in HPC systems architecture. In particular, a transition from a compute centric to a data movement centric point of view needs to be considered. Alternatives including quantum and neuromorphic computing have also been considered. The prospects of these technologies for post-Moore’s Law supercomputing will be explored.

MuMMI: A Modeling Infrastructure for Exploring Power and Execution Time Tradeoffs

Valerie Taylor (Texas A&M University)

Bio: Valerie Taylor is the Senior Associate Dean of Academic Affairs in the Dwight Look College of Engineering and the Regents Professor and Royce E. Wisenbaker Professor in the Department of Computer Science and Engineering at Texas A&M University. In 2003, she joined Texas A&M University as the Department Head of CSE, where she remained in that position until 2011. Prior to joining Texas A&M, Dr. Taylor was a member of the faculty in the EECS Department at Northwestern University for eleven years. She has authored or co-authored over 100 papers in the area high performance computing. She is also the Executive Director of the Center for Minorities and People with Disabilities in IT (CMD-IT). Dr. Taylor is an IEEE Fellow and has received numerous awards for distinguished research and leadership, including the 2001 IEEE Harriet B. Rigas Award for a woman with significant contributions in engineering education, the 2002 Outstanding Young Engineering Alumni from the University of California at Berkeley, the 2002 CRA Nico Habermann Award for increasing the diversity in computing, and the 2005 Tapia Achievement Award for Scientific Scholarship, Civic Science, and Diversifying Computing. Dr. Taylor is a member of ACM. Valerie E. Taylor earned her B.S. in ECE and M.S. in Computer Engineering from Purdue University in 1985 and 1986, respectively, and a Ph.D. in EECS from the University of California, Berkeley, in 1991.

Abstract: MuMMI (Multiple Metrics Modeling Infrastructure) is an infrastructure that facilitates the measurement and modeling of performance and power for large-scale paral-

lel applications. MuMMI builds upon three existing tools: Prophesy for performance modeling and prediction of parallel applications, PAPI for hardware performance counter monitoring, and PowerPack for power measurement and profiling. The MuMMI framework develops models of runtime and power based upon performance counters; these models are used to explore tradeoffs and identify methods for improving energy efficiency. In this talk we will describe the MuMMI framework and present examples of the use of MuMMI to improve the energy efficiency of parallel applications.

Wednesday, November 19

Invited Talks

Chair: Robert F. Lucas (Information Sciences Institute)

10:30am-12pm

Room: New Orleans Theater

A Curmudgeon's View of High Performance Computing

Michael Heath (University of Illinois at Urbana-Champaign)

Bio: Michael T. Heath is Professor and Fulton Watson Copp Chair Emeritus in the Department of Computer Science at the University of Illinois at Urbana-Champaign. He received a B.A. in Mathematics from the University of Kentucky, an M.S. in Mathematics from the University of Tennessee, and a Ph.D. in Computer Science from Stanford University. Before joining the University of Illinois in 1991, he spent a number of years at Oak Ridge National Laboratory, first as Eugene P. Wigner Postdoctoral Fellow and later as Computer Science Group Leader in the Computer Science and Mathematics Division. At Illinois, he served as Director of the Computational Science and Engineering Program 1996-2012, Director of the Center for Simulation of Advanced Rockets 1997-2010, and Interim Head of the Department of Computer Science 2007-2009. Heath's research interests are in scientific computing, particularly numerical linear algebra and optimization, and in parallel computing. He has been an editor of the SIAM Journal on Scientific Computing, SIAM Review, and the International Journal of High Performance Computing Applications, as well as several conference proceedings. He is also author of the widely adopted textbook *Scientific Computing: An Introductory Survey*, 2nd ed., published by McGraw-Hill in 2002. At the University of Illinois, Heath has been lead investigator on more than \$52M in research funding from federal and corporate grants and contracts. His major awards include being named an ACM Fellow by the Association for Computing Machinery in 2000, membership in the European Academy of Sciences in 2002, the Apple Award for Innovation in Science in 2007, the Taylor L. Booth Education Award from the IEEE Computer Society in 2009, and being named a SIAM Fellow by the Society for Industrial and Applied Mathematics and an Associate Fellow by the American Institute of Aeronautics and Astronautics, both in 2010.

Abstract: This talk addresses a number of concepts and issues in high performance computing that are often misconceived, misconstrued, or misused. Specific questions considered include why Amdahl remains relevant, what Moore actually predicted, the uses and abuses of speedup, and what should scalability really mean. These observations are based on more than thirty years experience in high performance computing, wherein one becomes an expert on mistakes by making a lot of them.

Supercomputing Trends, Opportunities and Challenges for the Next Decade

Jim Sexton (IBM)

Bio: Dr. James Sexton is the Program Director for the Computational Science Center and a senior manager at IBM T. J. Watson Research Center where he leads application co-design activities for IBM Data Centric Systems Department. Dr. Sexton received his Ph.D. in Theoretical Physics from Columbia University, NY. His areas of interest lie in High Performance Computing, Computational Science, Applied Mathematics and Analytics. Prior to joining IBM, Dr. Sexton held appointments as Lecturer then Professor at Trinity College Dublin, as postdoctoral fellow at IBM T. J. Watson Research Center, and held appointments at the Institute for Advanced Study at Princeton and at Fermi National Accelerator Laboratory. He has held adjunct appointments as Director and Founder of the Trinity Center for High Performance Computing, as a Board Member for the Board of Trinity College Dublin, as Senior Research Consultant for Hitachi Dublin Laboratory, and as a Hitachi Research Fellow at Hitachi's Central Research Laboratory in Tokyo. Dr. Sexton has over 70 publications and has participated on three separate Gordon Bell Award winning teams.

Abstract: The last decade has been an extraordinary period in the development of supercomputing. Computing performance has continued to scale with Moore's law even through the end of classic silicon scaling, 100+ Petaflop systems are on the horizon, architectural innovations abound and commercial installations have grown to rival traditional research and academic systems in capability and performance. The next decade promises to be even more extraordinary. Exascale will soon be in reach---architectures are embracing heterogeneity at unprecedented levels, commercial uses are driving capability in unexpected directions and fundamentally new technologies are beginning to show promise. In this presentation, we will discuss some of the emerging trends in supercomputing, some of the emerging opportunities both for technical and commercial uses of supercomputing, and some of the very significant challenges that are increasingly threatening to inhibit progress. The next decade will be a great time to be a systems designer!

Invited Talks

Chair: William Kramer (National Center for Supercomputing Applications)

3:30pm-5pm

Room: New Orleans Theater

Using Supercomputers to Discover the 100 Trillion Bacteria Living Within Each of Us

Larry Smarr (University of California, San Diego)

Bio: Larry Smarr is the founding Director of the California Institute for Telecommunications and Information Technology (Calit2), a UC San Diego/UC Irvine partnership, and holds the Harry E. Gruber professorship in the Department of Computer Science and Engineering (CSE) of UCSD's Jacobs School of Engineering. Before that he was the founding director of the National Center for Supercomputing Applications (NCSA) at the University of Illinois at Champaign-Urbana. He is a member of the National Academy of Engineering, as well as a Fellow of the American Physical Society and the American Academy of Arts and Sciences. In 2006 he received the IEEE Computer Society Tsutomu Kanai Award for his lifetime achievements in distributed computing systems. He is a member of the DOE ESnet Policy Board. He served on the NASA Advisory Council to 4 NASA Administrators, was chair of the NSF Advisory Committee on Cyberinfrastructure for the last 3 years, and for 8 years he was a member of the NIH Advisory Committee to the NIH Director, serving 3 directors. He was PI of the NSF OptIPuter project and of the Moore Foundation CAMERA global microbial metagenomics computational repository. His personal interests include growing orchids, snorkeling coral reefs, and quantifying the state of his body. You can follow him on his life-streaming portal at <http://lsmarr.calit2.net>.

Abstract: The human body is host to 100 trillion microorganisms, ten times the number of cells in the human body, and these microbes contain 100 times the number of DNA genes our human DNA does. The microbial component of our "super-organism" is comprised of hundreds of species with immense biodiversity. Thanks to the National Institutes of Health's Human Microbiome Program researchers have been discovering the states of the human microbiome in health and disease. To put a more personal face on the "patient of the future," I have been collecting massive amounts of data from my own body over the last five years, which reveals detailed examples of the episodic evolution of this coupled immune-microbial system. An elaborate software pipeline, running on high performance computers, reveals the details of the microbial ecology and its genetic components. We can look forward to revolutionary changes in medical practice over the next decade.

MATLAB Meets the World of Supercomputing

Cleve Moler (MathWorks)

Bio: Cleve Moler is chief mathematician, chairman, and cofounder of MathWorks. Moler was a professor of math and computer science for almost 20 years at the University of Michigan, Stanford University, and the University of New Mexico. He spent five years with two computer hardware manufacturers, the Intel Hypercube organization and Ardent Computer, before joining MathWorks full-time in 1989. In addition to being the author of the first version of MATLAB, Moler is one of the authors of the LINPACK and EISPACK scientific subroutine libraries. He is coauthor of three traditional textbooks on numerical methods and author of two online books, Numerical Computing with MATLAB and Experiments with MATLAB.

Abstract: MATLAB played a role in the very first Supercomputing conferences in 1988 and 1989. But then, for 15 years, MathWorks had little to do with the world of supercomputing. We returned in 2004 to SC04 with the introduction of parallel computing capability in MATLAB. We have had an increasing presence at the SC conferences ever since. I will review our experience over these last 10 years with tools for high performance computing ranging from GPUs and multicore to clusters and clouds.

Thursday, November 20

Plenary Invited Talks

Chair: Padma Raghavan (The Pennsylvania State University)

8:30am-10am

Room: New Orleans Theater

Meeting the Computational Challenges Associated with Human Health

Philip Bourne (National Institutes of Health)

Bio: Philip E. Bourne PhD is the Associate Director for Data Science (ADDS) at the National Institutes of Health. Formally he was Associate Vice Chancellor for Innovation and Industry Alliances, a Professor in the Department of Pharmacology and Skaggs School of Pharmacy and Pharmaceutical Sciences at the University of California San Diego, Associate Director of the RCSB Protein Data Bank and an Adjunct Professor at the Sanford Burnham Institute. Bourne's professional interests focus on service and research. He serves the national biomedical community through contributing ways to maximize the value (and hence accessibility) of scientific data. His research focuses on relevant biological and educational outcomes derived from computation and scholarly communication. This implies algorithms, text mining, machine learning, metalanguages, biological databases, and visualization applied to problems in systems pharmacology, evolution, cell signaling, apoptosis, immunology and scientific dissemination. He has published

over 300 papers and 5 books, one of which sold over 150,000 copies. Bourne is committed to maximizing the societal benefit derived from university research. Previously he co-founded four companies: ViSoft Inc., Protein Vision Inc. (a company distributing independent films for free), and most recently SciVee. Bourne is committed to furthering the free dissemination of science through new models of publishing and better integration and subsequent dissemination of data and results which as far as possible should be freely available to all. He is the co-founder and founding Editor-in-Chief of the open access journal PLOS Computational Biology. Bourne is committed to professional development through the Ten Simple Rules series of articles and a variety of lectures and video presentations. Bourne is a Past President of the International Society for Computational Biology, an elected fellow of the American Association for the Advancement of Science (AAAS), the International Society for Computational Biology (ISCB) and the American Medical Informatics Association (AMIA).

Abstract: It is my guess that by the end of this decade healthcare will be a predominantly digital enterprise and for the first time (surprising to say perhaps) patient centric. Such a shift from analog research and diagnosis and a provider centric healthcare system to an open digital system is a major change with opportunities and challenges across all areas of computation. I will describe some of these challenges and opportunities and indicate how the NIH is meeting them and how we need your help.

The Transformative Impact of Parallel Computing for Real-Time Animation

Lincoln Wallen (Dreamworks Animation)

Bio: Lincoln Wallen serves as Chief Technology Officer for DreamWorks Animation and is responsible for providing strategic technology vision and leadership for the Company. He joined DreamWorks Animation in 2008 as Head of Research and Development, where he was oversaw the strategic vision, creation and deployment of the studio's CG production platform and software tools. Wallen is a recipient of InfoWorld's 2012 "Technology Leadership Awards," the prestigious annual award that honors senior executives who have demonstrated creative, effective leadership in inventing, managing, or deploying technology within their organizations or in the IT community. Under his leadership, DreamWorks Animation was named to the 50 Most Innovative Companies list in MIT's *Technology Review*. Prior to joining DreamWorks Animation, Wallen was Chief Technology Officer for Online and Mobile Services at Electronic Arts where he led the gaming company's technical approach to publishing and videogame development for cell phones. He also served as Vice President at Criterion Software, Ltd, and as Chief Technology Officer for MathEngine. Before joining the entertainment industry, Wallen enjoyed a distinguished career as a Professor at Oxford University and was the first Director for the multidisciplinary Smith Institute

for Industrial Mathematics and Systems Engineering. He holds a B.S. in Mathematics and Physics from the University of Durham, United Kingdom, and a PhD in Artificial Intelligence as well as an honorary Doctorate from the University of Edinburgh, United Kingdom.

Abstract: Over the course of 20 years, DreamWorks Animation has been a leader in producing award winning computer generated (CG) animated movies such as Shrek, Madagascar and How to Train Your Dragon 2. In order to scale numerous simulations such as explosions or hundreds of animated characters in a scene, the studio requires an HPC environment capable of delivering tens of millions of render hours and managing millions of files for one film. Five years ago, the studio made a decision to re-architect key tools for animation and rendering from the ground up. Through the extensive use of parallel computing the studio has solved the fundamental need to maximize artist productivity and work at the speed of creativity. DreamWorks Animation demonstrates the fully transformative impact ubiquitous parallel computing will have on the complex and varied workflows in animation.

Invited Talks

Chair: Padma Raghavan (The Pennsylvania State University)

10:30am-12pm

Room: New Orleans Theater

Life at the Leading Edge:

Supercomputing @ Lawrence Livermore

Dona Crawford (Lawrence Livermore National Laboratory)

Bio: As Associate Director for Computation at Lawrence Livermore National Laboratory, Dona L. Crawford leads the laboratory's high performance computing efforts. This includes managing a staff who develop and deploy an integrated computing environment for petascale simulations of complex physical phenomena such as understanding global climate warming, clean energy creation, biodefense, and non-proliferation. This environment includes high performance computers, scientific visualization facilities, high-performance storage systems, network connectivity, multi-resolution data analysis, mathematical models, scalable numerical algorithms, computer applications, and necessary services to enable Laboratory mission goals and scientific discovery through simulation. An icon for the computing environment provided is the Advanced Simulation and Computing (ASC) Program's BlueGene/Q Sequoia machine (peak 20 PF). This is among the fastest computers in the world. Ms. Crawford has served on advisory committees for the National Research Council and the National Science Foundation. She is Co-Chair of the CRDF Global Board, Co-Chair of the Council on Competitiveness High Performance Computing Advisory Committee, and a member of IBM's Deep Computing Institute's External Advisory Board. She is a member of the California Council on Science and Technology, the IEEE and the ACM. In November 2010, Dona was featured

as one of insideHPC's "Rock Stars of HPC." Her undergraduate alma mater, presented her with the "Alumni Career Achievement Award" in 2012, and HPCwire named her among their 2013 People to Watch.

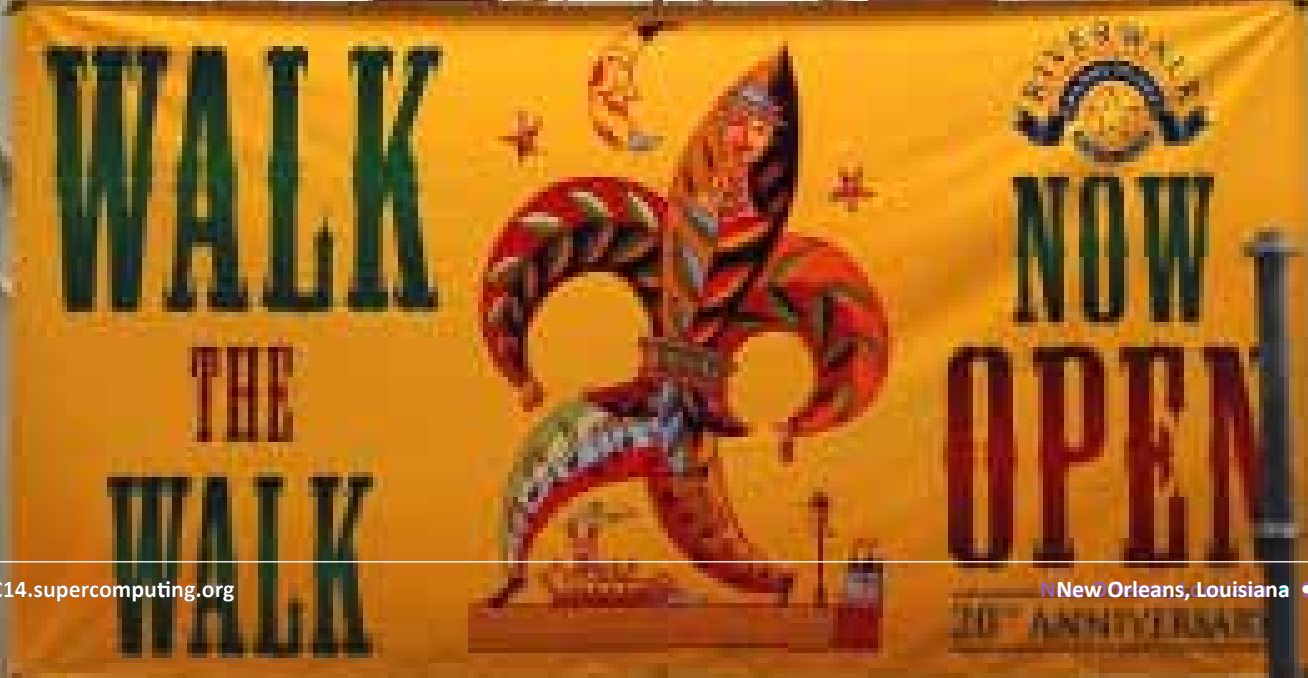
Abstract: Even before Livermore Laboratory opened in 1952, Ernest O. Lawrence and Edward Teller placed an order for one of the first commercial supercomputers, a Univac mainframe. Although computer simulation was in its infancy at the time, Laboratory leaders recognized its potential to accelerate scientific discovery and engineering. The Lab focused on complex, long-term national security problems that could be advanced through computing. These mission-focused applications drove the required computing hardware advances, which drove the associated infrastructure. What factors allowed LLNL to become leaders in HPC? How have they been sustained over the years? Come hear the inside Lab perspective on why #hpcmatters and what the Lab is doing to advance the discipline.

Runtime Aware Architectures

Mateo Valero (Barcelona Supercomputing Center)

Bio: Mateo Valero is a professor in the Computer Architecture Department at UPC in Barcelona. His research interests focus on high performance architectures. He has published approximately 600 papers, has served in the organization of more than 300 International Conferences and has given more than 400 invited talks. He is the director of the Barcelona Supercomputing Centre, the National Centre of Supercomputing in Spain. Dr. Valero has been honored with several awards. Among them, the Eckert-Mauchly Award, Harry Goode Award, The ACM Distinguish Service award, The "King Jaime I" in research and two Spanish National Awards on Informatics and on Engineering. He has been named Honorary Doctor by the Universities of Chalmers, Belgrade and Veracruz in Mexico and by the Spanish Universities of Las Palmas de Gran Canaria, Zaragoza and Complutense in Madrid. "Hall of the Fame" member of the IST European Program (selected as one of the 25 most influential European researchers in IT during the period 1983-2008). Professor Valero is Academic member of the Royal Spanish Academy of Engineering, of the Royal Spanish Academy of Doctors, of the Academia Europaea, and of the Academy of Sciences in Mexico, and Correspondant Academic of the Spanish Royal Academy of Science. He is a Fellow of the IEEE, Fellow of the ACM and an Intel Distinguished Research Fellow.

Abstract: In the last few years, the traditional ways to keep the increase of hardware performance to the rate predicted by Moore's Law have vanished. When uni-cores were the norm, hardware design was decoupled from the software stack thanks to a well-defined Instruction Set Architecture (ISA). This simple interface allowed developing applications without worrying too much about the underlying hardware, while hardware designers were able to aggressively exploit instruction-level parallelism (ILP) in superscalar processors. Current multi-cores are designed as simple symmetric multiprocessors (SMP) on a chip. However, we believe this is not enough to overcome all the problems facing multi-cores. The runtime has to drive the design of future multi-cores to overcome the restrictions in terms of power, memory, programmability and resilience that multi-cores have. In this talk, we introduce a first approach towards a Runtime-Aware Architecture (RAA), a massively parallel architecture designed from the runtime's perspective.



Panels

Panels at SC14 will be, as in past years, among the most important and heavily attended events of the conference. Panels will bring together the key thinkers and producers in the field to consider in a lively and rapid-fire context some of the key questions challenging high-performance computing, networking, storage and associated analysis technologies for the foreseeable future.

Panels bring a rare opportunity for mutual engagement of community leaders and broad mainstream contributors in a face-to-face exchange through audience participation and questioning. Surprises are the norm at panels, which makes for exciting and lively hour-and-a-half sessions. Panels can explore topics in-depth by capturing the opinions of a wide range of people active in the relevant fields. Panels represent state-of-the-art opinions, and can be augmented with social media technologies including Twitter, LinkedIn, and video feeds – even real-time audience polling. Please plan on actively participating in one or more of the panel offerings. We look forward to everyone's participation in making panels at SC14 a major success and lots of fun.

Panels

Tuesday, November 18

Funding Strategies for HPC Software Beyond Borders

10:30am-12pm

Room: 383-84-85

Moderator: Sabine Roller (University of Siegen)

Panelists: Marcus Wilms (German Research Foundation), Akinori Yonezawa (Riken Advanced Institute of Computational Sciences), Irene Qualters (National Science Foundation), Yutong Lu (National University of Defense Technology, China), Horst Simon (Lawrence Berkeley National Laboratory), Sophie Valcke (European Center for Research and Advanced Training in Scientific Computation)

HPC software for the exascale age becomes more and more sophisticated and complex. Challenges facing global issues like climate change, energy production, water supply, environmental impact, or natural disasters require international efforts and collaboration. This holds not only for politics and decision makers, but even more for the scientists providing insight into the global developments by simulation studies. Today, to enable a sustainable HPC software development, funding strategies need to go beyond borders. To support international and interdisciplinary collaboration, to ensure rapid developments of software and tools as soon as hardware progresses, several international funding programs have been established in the past on global (G8) as well as European or Asian-Pacific level in multi-, bi- or tri-lateral programs. The panel will discuss possibilities for future multi-lateral funding strategies for international collaboration, bringing together and supporting scientists all over the world.

Changing Operating Systems is Harder than Changing Programming Languages

1:30pm-3pm

Room: 383-84-85

Moderator: Arthur B. Maccabe (Oak Ridge National Laboratory)

Panelists: Marc Snir, Pete Beckman (Argonne National Laboratory), Patrick Bridges (University of New Mexico), Ron Brightwell (Sandia National Laboratories), Orran Krieger (Boston University), John Kubiawicz (University of California, Berkeley), Satoshi Matsuoka (Tokyo Institute of Technology), Hank Hoffman (University of Chicago)

Since its introduction in 1970, Unix (including Linux) has been a mainstay in scientific computing. The use of Linux has many advantages, but also is a drag on innovation: Since Linux is used on a broad range of platforms, changes that are specific to HPC are accepted very slowly, if at all, by Linux maintainers.

In addition, very little work has focused on the “global operating systems”, i.e., the software system infrastructure needed to manage the entire machine and support multi-node parallel applications. Current Linux software will be stretched in dealing with future nodes with hundreds of heterogeneous cores and heterogeneous memory; it will continue to provide little support to the integration of local OS'es into a global OS. A panel composed of HPC OS/R experts will consider the technical, research, and socio-economic challenges that need to be addressed for future, extreme-scale operating systems.

HPC Productivity or Performance: Choose One

3:30pm-5pm

Room: 383-84-85

Moderator: William Harrod (Department of Energy Office of Science)

Panelists: Michael Heroux (Sandia National Laboratories), Robert Wisnieski (Intel Corporation), William Gropp (University of Illinois), Thomas Sterling (Indiana University), Satoshi Matsuoka (Tokyo Institute of Technology)

HPC faces a dramatic challenge between achieving high performance for applications and the ease with which such codes can be developed. Productivity is a notional quality measure that interrelates the two along with associated costs to provide a single unifying metrics. Although presently ill defined, productivity is essential in delivering effective computing environments, especially as HPC moves towards the exascale era. This panel will highlight the tensions among performance, usability, and cost (including energy) of extreme-scale computing and discuss the implications of the tradeoffs to drive progress toward superior computing methods and capabilities. Experts across diverse fields will consider what is required to make future systems much easier to use and for a wider range of problems even as scaling increases by orders of magnitude. One key issue is legacy code and retaining their investment for future system classes. Audience engagement will be driven by questions directed to attendees.

Wednesday, November 19

Analyst Crossfire and Advanced Cyberinfrastructure

10:30am-12pm

Room: 383-84-85

Moderators: Addison Snell (Intersect360 Research)

Panelists: Bill Kramer (National Center for Supercomputing Applications), Gilad Shainer (Mellanox Technologies), Sumit Gupta (NVIDIA Corporation)

In this fast-paced, opinion-driven panel, Addison Snell of Intersect360 Research will grill industry expert panelists from both the technology vendor and supercomputing user community on the emerging topics and trends from SC14.

In this popular Intersect360 Research format, four panelists—two vendors and two users—will face off on topics selected by the analyst to be of particular interest for the road ahead. With a clock on the screen, answers will be direct, and designed to highlight the complexities of the technologies and trends developing in supercomputing. The panelists will complete seven topics in 30 minutes, including a short answer “speed round.” Audience Q&A will follow.

Advanced Cyberinfrastructure - Research and Education Facilitator Program

Moderator: James B. von Oehsen (Clemson University)

Panelists: Ken-ichi Nomura (University of Southern California), Galen Collier (Clemson University), Anita Orendt (University of Utah), Alex Feltus (Clemson University), Gwen Jacobs (University of Hawaii)

This session will describe current and future initiatives of the “Advanced Cyberinfrastructure Research and Education Facilitators (ACI-REFs)” program, a consortium that is forging a nationwide alliance of educators to empower local campus researchers to be more effective users of advanced cyberinfrastructure (ACI). The project engages the “long tail” of ACI users by building a coordinated network of ACI-REFs, computational scientists whose mission it is to leverage existing resources and “make a difference” in supporting their local campus researchers.

Panelists will present examples of successful strategies for supporting and working with researchers from disciplines that have not been traditional users of advanced cyberinfrastructure. The panel will bring together existing ACI-REFs, who will present ideas and best practices in how to engage these communities, and faculty researchers who have worked with the project and who will offer their perspective of the program and its potential to facilitate new knowledge discovery.

Can We Avoid Building An Exascale “Stunt” Machine?

1:30pm-3pm

Room: 383-84-85

Moderator: Mike Bernhardt (Intel Corporation)

Panelists: William Harrod (Department of Energy Office of Advanced Scientific Computing Research), Steve Scott (Cray Inc.), Alan Gara (Intel Corporation), Terri Quinn (Lawrence Livermore National Laboratory), Thomas Sterling (Indiana University), Rick Stevens (Argonne National Laboratory)

While most HPC stakeholders agree there is no value to building ‘stunt’ machines, there is a rolling thunder of opinion that an exascale “stunt” machine may be the smart thing to do. Many believe we – the community – should collaborate (vendors, users, funding agents) to build an exascale stunt machine. This system would be a proving ground or workbench, enabling us to test and prove aspects of functionality and capability from scaling to resiliency. Others believe it would be a terrible waste of money and effort, only delaying the more fruitful efforts of building a practical exascale-class system. This panel of industry opinion shapers will debate both sides of this discussion.

This panel of HPC leaders will debate the pros and cons of going down such a path as we explore what some scientists now believe is inevitable - if we are ever going to achieve practical exaFLOPS-level performance.

Future of Memory Technology for Exascale and Beyond

3:30pm-5pm

Room: 383-84-85

Moderator: Richard Murphy (Micron Technology, Inc.)

Panelists: Shekhar Borkar (Intel), Bill Dally (NVIDIA Corporation), Andreas Hansson (ARM Ltd.), Mike Ignatowski (Advanced Micro Devices, INC.), Doug Joseph (IBM Corporation), Peter Kogge (University of Notre Dame), Troy Manning (Micron Technology, Inc.)

Memory technology is in the midst of profound change as we move into the exascale era. Early analysis, including the DARPA UHPC Exascale Report correctly identified the fundamental technology problem as one of enabling low-energy data movement throughout the system. However, the end of Dennard Scaling and the corresponding impact on Moore’s Law has begun a fundamental transition in the relationship between the processor and memory system. The lag in the increase in the number of cores compared to what Moore’s Law would provide has proven a harbinger of the trend towards memory systems performance dominating compute capability. This panel began to address some of the critical exascale questions at SC13 and will continue the discussions at SC14.

Thursday, November 20

Multi-System Supercomputer Procurements – Is It a Good Idea?

10:30am-12pm

Room: 383-84-85

Moderator: Arthur S. Bland (Oak Ridge National Laboratory)

Panelists: Bronis R. de Supinski (Lawrence Livermore National Laboratory), Manuel Vigil (Los Alamos National Laboratory), Katie Antypas (Lawrence Berkeley National Laboratory), Susan Coghlan (Argonne National Laboratory), James H. Rogers (Oak Ridge National Laboratory)

The U.S. Department of Energy's national laboratories have partnered to issue two solicitations to purchase five major supercomputers for delivery over the next several years. Los Alamos, Sandia, and Lawrence Berkeley National Laboratories are procuring the Trinity and NERSC-8 systems and Lawrence Livermore, Argonne, and Oak Ridge National Laboratories partnered to procure the Sierra, ALCF-3, and OLCF-4 systems. In this panel, we will briefly describe the acquisitions and then discuss the pros and cons of issuing a single solicitation to acquire multiple systems. The moderator and audience will pose questions on the process: was this a good way to acquire systems or not?; do they think the outcome of the acquisition was better or not? The panel will include project or technical leads for each system, as well as representatives from the vendor community.

InfoSymbioticSystems/DDDAS – The Power of Dynamic Data-Driven Application Systems and the Next Generation of Big-Computing and Big-Data

1:30pm-3pm

Room: 383-84-85

Moderator: Frederica Darema (Air Force Office of Scientific Research)

Panelists: Kelvin Droegemeier (University of Oklahoma), Abani Patra (University at Buffalo, The State University of New York), Nurcin Celik (University of Miami), Vaidy Sunderam (Emory University), Subhasish Mitra (Stanford University), Christos Kozyrakis (Stanford University), Aniruddha Gokhale (Vanderbilt University)

Dynamic Data-Driven Application Systems (DDDAS) is a paradigm whereby application-simulation models of natural and engineered systems become a symbiotic feedback control system with the application's instrumentation measurements. Through this dynamic integration across computing and instrumentation DDDAS creates new capabilities for more accurate analysis, prediction, and control in application systems. The

instrumentation data considered, real-time or archival, and resulting from heterogeneous sensor- and actuation-networks and mobile devices, are the next wave of Big Data. DDDAS enables intelligent management of such Big Data and extends the traditional notions of "Big Computing" to encompass the diverse range of platforms from the exascale to sensors and controllers and to mobile systems. The panel will discuss new DDDAS-enabled and DDDAS-driven capabilities in important application areas such as civilian and national security infrastructures; environmental systems; medical care systems; privacy and security; systems software and hardware supporting DDDAS environments and in the context of commonalities in underlying exascale and sensor-scale technologies.

Challenges with Liquid Cooling - A Look into the Future of HPC Data Centers

3:30pm-5pm

Room: 383-84-85

Moderator: Torsten Wilde (Leibniz Supercomputing Center)

Panelists: Michael Patterson (Intel Corporation), Thomas Blum (Megware), Laurent Cargemel (Bull), Giovanbattista Mattiussi (Eurotech), Nicolas Dube (Hewlett-Packard)

Liquid cooling is key to dealing with the heat density, reducing energy consumption, and increasing the performance of this generation of supercomputers and becomes even more predominant on the roadmap for the foreseeable future. The transition to liquid cooling, however, comes with challenges.

One challenge is that each system and each site comes with its specific issues regarding liquid cooling. These complicate procurements, design, installation, operations, and maintenance.

This panel will review the immediate past history from a lessons learned perspective as well as discuss what's needed for liquid cooling to be implemented more readily in the future. This includes challenges such as data and metrics to assess the relative efficiencies of the different cooling technologies, water quality, heat recovery, the disparity between building life and cluster life, and other issues.

How can we as a community address these challenges?

Friday, November 21

Beyond Von Neumann, Neuromorphic Systems and Architectures

8:30am-10am

Room: 391-92

Moderator: Mark E. Dean (University of Tennessee, Knoxville)

Panelists: Daniel Hammerstrom (Defense Advanced Research Projects Agency), Jacob Vogelstein (National Intelligence Advanced Research Projects Activity), Karlheinz Meier (Heidelberg University), Jeff Krichmar (University of California, Irvine), Dhireesha Kudithipudi (Rochester Institute of Technology), Dharmendra Moda (IBM)

Technology has enabled people, businesses, governments, and societies to be more instrumented, interconnected and knowledgeable of their environment and socio-economic conditions that can impact their wellbeing. Unfortunately, conventional computing systems have a difficult time pulling timely and useful insight from the tsunami of real-time, noisy, low-precision, uncorrelated, spatio-temporal, multimodal, “natural” data that are collected. Conventional systems are constrained by power consumption, physical size, network limitations and processing speed. A new, unconventional approach is needed to move beyond the bottlenecks of von Neumann architecture systems to enable the efficient and effective use of exascale real-time natural data.

Panelists include representatives from academia, industry and government. The panel will discuss how recent progress in neuromorphic (neuro-inspired) architectures and systems will enable computational systems to gain insight from large real-time data-sets. The panelists will compare and contrast existing approaches, explore target applications, and consider the challenges faced in their deployment and adoption.

Return of HPC Survivor: Outwit, Outlast, Outcompute

8:30am-10am

Room: 383-84-85

Moderator: Cherri M. Pancake (Oregon State University)

Panelists: Cherri M. Pancake (Oregon State University), Mike Bernhardt (Intel Corp.), Keith Gray (BP), Pete Beckman (Argonne National Laboratory), John West (Department of Defense HPC Modernization Program), Dona Crawford (Lawrence Livermore National Laboratory), Phil Bourne (National Institutes of Health)

Back by popular demand, this panel brings together HPC experts to compete for the honor of “HPC Survivor 2014.” Following up on the popular Xtreme Architectures (2004), Xtreme Programming (2005), Xtreme Storage (2007), and Build Me an Xascale (2010) competitions, the theme for this year is “Does HPC Really Matter?”

The contest is a series of “rounds,” each posing a specific question about system design, philosophy, implementation, or use. After contestants answer, a distinguished commentator furnishes additional wisdom to help guide the audience. At the end of each round, the audience votes (applause, boos, etc.) to eliminate a contestant. The last contestant remaining wins.

Cherri Pancake returns as emcee of the event. Highlights include internationally renowned panelists, lively critical commentary, and exit interviews giving eliminated candidates an opportunity to explain why the audience is “wrong” to vote them out.

ROI from Academic Supercomputing

10:30am-12pm

Room: 391-92

Moderator: Greg Newby (King Abdullah University of Science and Technology)

Panelists: David Lifka (Cornell University), Susan Fratkin (Coalition for Academic Scientific Computation), Amy Apon (Clemson University), Craig Stewart (Indiana University), Nicholas Berente (University of Georgia), Rudolf Eigenmann (National Science Foundation)

Return on Investment or ROI is a fundamental measure of effectiveness in business. In this panel, we will share approaches to assessing ROI for academic supercomputing. The panel will address the challenge that “returns” from supercomputing and other computationally-based research activities are often not financial. This is major distinction from other industrial sectors, where product sales, inventions, and patents might form the basis of ROI calculations. How should ROI be assessed for high performance computing in academic environments? What inroads to ROI calculations are underway by the panelists? What are challenges of ROI? Financial outcomes at colleges and universities, notably the receipt of research funding or contracts. However, even when there are financial outcomes, it can be difficult to apportion the revenues due to supercomputing versus other things. Non-financial outcomes are major intellectual products of colleges and universities, and this is therefore a departure from ROI calculations in industry.

The Roadblock Ahead – Power Usage for Storage and I/O

10:30am-12pm

Room: 383-84-85

Moderator: Roger Ronald (*System Fabric Works*)

Panelists: John Shalf (*Lawrence Berkeley National Laboratory*), Bill Boas (*Cray Inc.*), Kurt Keville (*MIT*), Ronald P. Luijten (*IBM Zurich Research Laboratory*), Jim Ang (*Sandia National Laboratories*)

Predicted energy expenditures for traditional high-end computing systems in 2018 indicate that instruction execution will be a surprisingly small percentage (<20%) of power usage for the entire system. Energy used for memory, network, and storage accesses will contribute almost 60% of the total. As subsequent advances and scaling in computing occur, the percentage of energy needed to move and store data appears likely to grow even further.

Moving to exascale requires innovative solutions for reducing memory and I/O power utilization.

In the past, access delays to external memory have been addressed by caching within processors. Are there equivalent options on the horizon that could address storage and I/O power usage to enable continued advances in leading-edge computing?

Papers

The SC14 Technical Papers program received 394 submissions covering a wide variety of research topics in HPC. We followed a rigorous peer review process with an author rebuttal period, careful management of conflicts, and at least four reviews per submission (as many as six reviews in several cases). At a two-day face-to-face committee meeting on June 16-17 in Houston, technical paper committee members discussed the submissions and finalized their decisions. At the conclusion of the meeting, the committee had accepted 84 papers, which corresponds to an acceptance rate of 21.3 percent. Fourteen of the accepted papers have been selected as finalists for the Best Paper (BP) and Best Student Paper (BSP) awards. The BP and BSP winners will be announced during the Awards ceremony.

Papers

Tuesday, November 18

Heterogeneity and Scaling in Applications

Chair: Justin Luitjens (NVIDIA)

10:30am-12pm

Room: 393-94-95

Lattice QCD with Domain Decomposition on Intel(R) Xeon Phi(TM) Co-Processors

Simon Heybrock (University of Regensburg), Balint Joo (Thomas Jefferson National Accelerator Facility), Dhiraj D. Kalamkar, Mikhail Smelyanskiy, Karthikeyan Vaidyanathan (Intel Corporation), Tilo Wettig (University of Regensburg), Pradeep Dubey (Intel Corporation)

The gap between the cost of moving data and the cost of computing continues to grow, making it ever harder to design iterative solvers on extreme-scale architectures. This problem can be alleviated by alternative algorithms that reduce the amount of data movement. We investigate this in the context of Lattice Quantum Chromodynamics and implement such an alternative solver algorithm, based on domain decomposition, on Intel(R) Xeon Phi(TM) co-processor (KNC) clusters. We demonstrate close-to-linear on-chip scaling to all 60 cores of the KNC. With a mix of single- and half-precision the domain-decomposition method sustains 400-500 Gflop/s per chip. Compared to an optimized KNC implementation of a standard solver [1], our full multi-node domain-decomposition solver strong-scales to more nodes and reduces the time-to-solution by a factor of 5.

Mapping to Irregular Torus Topologies and Other Techniques for Petascale Biomolecular Simulation

James C. Phillips, Yanhua Sun, Nikhil Jain, Eric J. Bohm, Laxmikant V. Kale (University of Illinois at Urbana-Champaign)

Currently deployed petascale supercomputers typically use toroidal network topologies in three or more dimensions. While these networks perform well for topology-agnostic codes on a few thousand nodes, leadership machines with 20,000 nodes require topology awareness to avoid network contention on communication-intensive codes. Topology adaptation is complicated by irregular node allocation shapes and holes due to dedicated input/output nodes or hardware failure. In the context of the popular molecular dynamics program NAMD, we present methods for mapping a periodic 3-D grid of fixed-size spatial decomposition domains to 3-D Cray Gemini and 5-D IBM Blue Gene/Q toroidal networks to enable hundred-million atom full machine simulations, and to similarly partition node allocations into compact domains for smaller simulations using multiple-copy algorithms. Additional enabling techniques are discussed, and performance is reported for NCSA Blue Waters, ORNL Titan, ANL Mira, TACC Stampede, and NERSC Edison.

A Volume Integral Equation Stokes Solver for Problems with Variable Coefficients

Dhairya Malhotra, Amir Gholami, George Biros (University of Texas at Austin)

We present a novel numerical scheme for solving the Stokes equation with variable coefficients in the unit box. Our scheme is based on a volume integral equation formulation. We employ a novel adaptive fast multipole method for volume integrals to obtain a scheme that is algorithmically optimal. To increase performance, we have integrated our code with both NVIDIA and Intel accelerators. Overall, our scheme supports non-uniform discretizations and is spectrally accurate. In our largest run we solved a problem with 20 billion unknowns, using a 14-order approximation for the velocity, on 2048 nodes of the Stampede platform at the Texas Advanced Computing Center. We achieved 0.656 PetaFLOPS for the overall code (20% efficiency) and 1 PetaFLOPS for the fast multipole part (40% efficiency). As an application example, we simulate Stokes flow in a porous medium with highly complex pore structure using a penalty formulation to enforce the no slip condition.

Award: Best Student Paper Finalist

Memory and Microarchitecture

Chair: Suzanne Rivoire (Sonoma State University)

10:30am-12pm

Room: 391-92

Fence Scoping

Changhui Lin (University of California, Riverside), Vijay Nagarajan (University of Edinburgh), Rajiv Gupta (University of California, Riverside)

We observe that fence instructions used by programmers are usually only intended to order memory accesses within a limited scope. Based on this observation, we propose the concept fence scope which defines the scope within which a fence enforces the order of memory accesses, called scoped fence (S-Fence). S-Fence is a customizable fence, which enables programmers to express ordering demands by specifying the scope of fences when they only want to order part of memory accesses. At runtime, hardware uses the scope information conveyed by programmers to execute fence instructions in a manner that imposes fewer memory ordering constraints than a traditional fence, and hence improves program performance. Our experimental results show that the benefit of S-Fence hinges on the characteristics of applications and hardware parameters. A group of lock-free algorithms achieve peak speedups ranging from 1.13x to 1.34x; while full applications achieve speedups ranging from 1.04x to 1.23x.

Recycled Error Bits: Energy-Efficient Architectural Support for Floating Point Accuracy

Ralph Nathan (Duke University), Bryan Anthonio (Cornell University), Shih-Lien L. Lu (Intel Corporation), Helia Naeimi (Intel Corporation), Daniel J. Sorin, Xiaobai Sun (Duke University)

In this work, we provide energy-efficient architectural support for floating point accuracy. For each floating point addition performed, we “recycle” that operation’s rounding error. We make this error architecturally visible such that it can be used, whenever desired, by software. We also design a compiler pass that allows software to automatically use this feature. Experimental results on physical hardware show that software that exploits architecturally recycled error bits can (a) achieve accuracy comparable to a 64-bit FPU with performance and energy that are comparable to a 32-bit FPU, and (b) achieve accuracy comparable to an all-software scheme for 128-bit accuracy with far better performance and energy usage.

Managing DRAM Latency Divergence in Irregular GPGPU Applications

Niladrish Chatterjee (NVIDIA / University of Utah), Mike O’Connor (NVIDIA / University of Texas at Austin), Gabriel H. Loh (Advanced Micro Devices, Inc.), Nuwan Jayasena (Advanced Micro Devices, Inc.), Rajeev Balasubramonian (University of Utah)

Memory controllers in modern GPUs aggressively reorder requests for high bandwidth usage, often interleaving requests from different warps. This leads to high variance in the latency of different requests issued by the threads of a warp. Since a warp in a SIMT architecture can proceed only when all of its memory requests are returned by memory, such latency divergence causes significant slowdown when running irregular GPGPU applications.

To solve this issue, we propose memory scheduling mechanisms that avoid inter-warp interference in the DRAM system to reduce the average memory stall latency experienced by warps. We further reduce latency divergence through mechanisms that coordinate scheduling decisions across multiple independent memory channels. Finally we show that carefully orchestrating the memory scheduling policy can achieve low average latency for warps, without compromising bandwidth utilization. Our combined scheme yields a 10.1% performance improvement for irregular GPGPU workloads relative to a throughput-optimized GPU memory controller.

Performance Measurement

Chair: Mitsuhsa Sato (University of Tsukuba)

10:30am-12pm

Room: 388-89-90

CYPRESS: Combining Static and Dynamic Analysis for Top-Down Communication Trace Compression

Jidong Zhai (Tsinghua University), Jianfei Hu (University of California, Irvine), Xiongchao Tang (Tsinghua University), Xiaosong Ma (Qatar Computing Research Institute), Wenguang Chen (Tsinghua University)

Communication traces are increasingly important, both for parallel applications’ performance analysis/optimization, and for designing next-generation HPC systems. Meanwhile, the problem size and the execution scale on supercomputers keep growing, producing prohibitive volume of communication traces. To reduce the size of communication traces, existing dynamic compression methods introduce large compression overhead with the job scale.

We propose a hybrid static-dynamic method that leverages information acquired from static analysis to facilitate more effective and efficient dynamic trace compression. Our proposed scheme, CYPRESS, extracts a program communication structure tree at compile time using inter-procedural analysis. This tree naturally contains crucial iterative computing features such as the loop structure, allowing subsequent runtime compression to “fill in”, in a “top-down” manner, event details into the known communication template. Results show that CYPRESS reduces intra-process and inter-process compression overhead up to 5x and 9x respectively over state-of-the-art dynamic methods, while only introducing very low compiling overhead.

Award: Best Paper Finalist

Lightweight Distributed Metric Service: A Scalable Infrastructure for Continuous Monitoring of Large Scale Computing Systems and Applications

Anthony Agelastos, Benjamin Allan, Jim Brandt (Sandia National Laboratories); Paul Cassella (Cray Inc.); Jeremy Enos, Joshi Fullop (National Center for Supercomputing Applications); Ann Gentile, Steve Monk (Sandia National Laboratories); Nichamon Naksinehaboon (Open Grid Computing); Jeff Ogden, Mahesh Rajan (Sandia National Laboratories), Michael Showerman (National Center for Supercomputing Applications), Joel Stevenson (Sandia National Laboratories); Narate Taerat, Tom Tucker (Open Grid Computing)

Understanding how resources of High Performance Compute platforms are utilized by applications both individually and as a composite is key to application and platform performance. Typical system monitoring tools do not provide sufficient

fidelity while application profiling tools do not capture the complex interplay between applications competing for shared resources. To gain new insights, monitoring tools must run continuously, system wide, at frequencies appropriate to the metrics of interest while having minimal impact on application performance.

We introduce the Lightweight Distributed Metric Service for scalable, lightweight monitoring of large scale computing systems and applications. We describe issues and constraints guiding deployment in Sandia National Laboratories' capacity computing environment and on the National Center for Supercomputing Applications' Blue Waters platform including motivations, metrics of choice, and requirements relating to the scale and specialized nature of Blue Waters. We address monitoring overhead and impact on application performance and provide illustrative profiling results.

Dissecting On-node Memory Access Performance: A Semantic Approach

Alfredo Gimenez (University of California, Davis), Todd Gamblin, Barry Rountree, Abhinav Bhatele (Lawrence Livermore National Laboratory); Ilir Jusufi (University of California, Davis), Peer-Timo Bremer (Lawrence Livermore National Laboratory), Bernd Hamann (University of California, Davis)

Optimizing memory access is critical for performance and power efficiency. CPU manufacturers have developed sampling-based performance measurement units (PMUs) that report precise costs of memory accesses at specific addresses. However, this data is too low-level to be meaningfully interpreted and contains an excessive amount of irrelevant or uninteresting information.

We have developed a method to gather fine-grained memory access performance data for specific data objects and regions of code with low overhead and attribute semantic information to the sampled memory accesses. This information provides the context necessary to more effectively interpret the data. We have developed a tool that performs this sampling and attribution and used the tool to discover and diagnose performance problems in real-world applications. Our techniques provide useful insight into the memory behavior of applications and allow programmers to understand the performance ramifications of key design decisions: domain decomposition, multi-threading, and data motion within distributed memory systems.

Accelerators

1:30pm-3pm

Room: 388-89-90

Practical Symbolic Race Checking of GPU Programs

Peng Li (University of Utah), Guodong Li (Fujitsu Labs of America), Ganesh Gopalakrishnan (University of Utah)

Even the careful GPU programmer can inadvertently introduce data races while writing and optimizing code. Currently available GPU race checking methods fall short either in terms of their formal guarantees, ease of use, or practicality. Existing symbolic methods: (1) do not fully support existing CUDA kernels; (2) may require user-specified assertions or invariants; (3) often require users to guess which inputs may be safely made concrete; (4) tend to explode in complexity when the number of threads is increased; and (5) explode in the face of thread-ID based decisions, especially in a loop. We present SESA, a new tool combining Symbolic Execution and Static Analysis to analyze C++ CUDA programs that overcomes all these limitations. SESA also scales well to handle non-trivial benchmarks such as Parboil and Lonestar, and is the only tool of its class that handles such practical examples. This paper presents SESA's methodological innovations and practical results.

Scalable Kernel Fusion for Memory-Bound GPU Applications

Mohamed Wahib, Naoya Maruyama (RIKEN Advanced Institute for Computational Science)

GPU implementations of HPC applications relying on finite difference methods can include tens of kernels that are memory-bound. Kernel fusion can improve the performance by reducing data traffic to off-chip memory; kernels that share data arrays are fused to larger kernels where on-chip cache is used to hold the data reused by instructions originating from different kernels. The main challenges are: a) Searching for the optimal kernel fusions while constrained by data dependences and kernels' precedences and, b) Effectively applying kernel fusion to achieve speedup. This paper introduces a problem definition and a scalable method for searching the space of possible kernel fusions to identify optimal kernel fusions for large problem sizes. The paper also introduces a codeless performance upper-bound projection to achieve effective fusions. Results show how using the proposed kernel fusion method improved the performance of two real-world applications containing tens of kernels by 1.35x and 1.2x.

A Unified Programming Model for Intra- and Inter-Node Offloading on Xeon Phi Clusters

Matthias Noack (Zuse Institute Berlin), Florian Wende (Zuse Institute Berlin), Frank Cordes (getLig&tar GbR), Thomas Steinke (Zuse Institute Berlin)

Standard offload programming models for the Xeon Phi, e.g. Intel LEO and OpenMP 4.0, are restricted to a single compute node and hence a limited number of coprocessors. Scaling applications across a Xeon Phi cluster/supercomputer thus requires hybrid programming approaches, usually MPI+X. In this work, we present a framework based on heterogeneous active messages (HAM-Offload) that provides the means to offload work to local and remote (co)processors using a unified offload API. Since HAM-Offload provides similar primitives as current local offload frameworks, existing applications can be easily ported to overcome the single-node limitation while keeping the convenient offload programming model. We demonstrate the effectiveness of the framework by using it to enable a real-world application from the field of molecular dynamics to use multiple local and remote Xeon Phis. The evaluation shows good scaling behavior. Compared with LEO, performance is equal for large offloads and significantly better for small offloads.

Best Practices in File Systems

Chair: Mark Gary (Lawrence Livermore National Laboratory)

1:30pm-3pm

Room: 393-94-95

Best Practices and Lessons Learned from Deploying and Operating Large-Scale Data-Centric Parallel File Systems

Sarp Oral, James Simmons, Jason Hill, Dustin Leverman, Feiyi Wang, Matthew Ezell, Ross Miller, Douglas Fuller, Raghul Gunasekaran, Youngjae Kim, Saurabh Gupta, Devesh Tiwari, Sudharshan Vazhkudai, James Rogers, Arthur S. Bland, Galen M. Shipman (Oak Ridge National Laboratory), David Dillow (Self Employed)

Oak Ridge Leadership Computing Facility (OLCF) has deployed multiple world-class parallel file systems to support its compute, data analysis, and visualization platforms. After deployment, OLCF has continued to hone operating strategies for these systems in response to rapidly changing technologies and user demands. OLCF's file systems have also served as test platforms for new file system features; as technology evaluation platforms for I/O strategies and benchmarks; and as hubs for data storage and transfer for OLCF users. During this process, OLCF has acquired significant expertise in the areas of data storage systems, file system software, technology evaluation, benchmarking, and procurement practices. This paper provides an account of our experience and lessons learned in

acquiring, deploying, and operating large, parallel file systems. We believe that these lessons will be useful to the wider HPC community involved in such activities.

Award: Best Paper Finalists

A User-Friendly Approach for Tuning Parallel File Operations

Robert McLay, Doug James, Si Liu, John Cazes, William Barth (University of Texas at Austin)

The Lustre file system provides high aggregated I/O bandwidth and is in widespread use throughout the HPC community. Here we report on work (1) developing a model for understanding collective parallel MPI write operations on Lustre, and (2) producing a library that optimizes parallel write performance in a user-friendly way. We note that a system's default stripe count is rarely a good choice for parallel I/O, and that performance depends on a delicate balance between the number of stripes and the actual (not requested) number of collective writers. Unfortunate combinations of these parameters may degrade performance considerably. For the programmer, however, it's all about the stripe count: an informed choice of this single parameter allows MPI to assign writers in a way that achieves near-optimal performance. We offer recommendations for those who wish to tune performance manually and describe the easy-to-use T3PIO library that manages the tuning automatically.

IndexFS: Scaling File System Metadata Performance with Stateless Caching and Bulk Insertion

Kai Ren, Qing Zheng, Swapnil Patil, Garth Gibson (Carnegie Mellon University)

The growing size of modern storage systems is expected to soon achieve and exceed billions of objects, making metadata operation critical to the overall performance. Many existing parallel and cluster file systems only focus on providing highly parallel access to file data, but lack a scalable metadata service. In this paper, we introduce a middleware design called IndexFS that adds support to existing file systems such as HDFS and PVFS for high-performance operations on metadata and small files. IndexFS uses a tabular-based architecture that incrementally partitions the namespace at per-directory basis, preserving disk locality for small directories. We also propose two client caching techniques: bulk insertion for creation intensive workloads and stateless metadata caching for hot spot mitigation. By combining these techniques, we scaled IndexFS to 128 servers for various metadata workloads. Experiments demonstrate that its out-of-core metadata throughput outperforms PVFS by 50% to an order of magnitude.

Award: Best Student Paper Finalist, Best Paper Finalist

Earth and Space Sciences

Chair: Brian O'Shea (Michigan State University)

1:30pm-3pm

Room: 391-92

High Productivity Framework on GPU-Rich Supercomputers for Operational Weather Prediction Code ASUCA

Takashi Shimokawabe, Takayuki Aoki, Naoyuki Onodera (Tokyo Institute of Technology)

The weather prediction code demands large computational performance to achieve fast and high-resolution simulations. Skillful programming techniques are required for obtaining good parallel efficiency on GPU supercomputers. Our framework-based weather prediction code ASUCA has achieved good scalability with hiding complicated implementation and optimizations required for distributed GPUs, contributing to increasing the maintainability; ASUCA is a next-generation high-resolution meso-scale atmospheric model being developed by the Japan Meteorological Agency. Our framework automatically translates user-written stencil functions that update grid points and generates both GPU and CPU codes. User-written codes are parallelized by MPI with intra-node GPU peer-to-peer direct access. These codes can easily utilize optimizations such as overlapping technique to hide communication overhead by computation. Our simulations on the GPU-rich supercomputer TSUBAME 2.5 at the Tokyo Institute of Technology have demonstrated good strong and weak scalability achieving 209.6 TFlops in single precision for our largest model using 4,108 NVIDIA K20X GPUs.

Pipelining Computational Stages of the Tomographic Reconstructor for Multi-Object Adaptive Optics on a Multi-GPU System

Ali Charara, Hatem Ltaief (King Abdullah University of Science & Technology), Damien Gratadour (Pierre and Marie Curie University), David Keyes (King Abdullah University of Science & Technology), Arnaud Sevin (Pierre and Marie Curie University), Ahmad Abdelfattah (King Abdullah University of Science & Technology), Eric Gendron, Carine Morel, Fabrice Vidal (Laboratory for Space Studies and Astrophysics Instrumentation, Paris Observatory, French National Center for Scientific Research, Pierre and Marie Curie University)

The European Extremely Large Telescope project (E-ELT) is one of the European highest priorities in ground-based astronomy. ELTs are built on top of a multitude of highly sensitive and critical astronomic instruments. In particular, a new instrument called MOSAIC has been proposed to perform multi-object spectrograph using a Multi-Object Adaptive Optics (MOAO) technique. The core implementation of the methodology lies in the intensive computational simulation of a tomographic

reconstructor, which accordingly adjusts, in return, the overall instrument for subsequent real-time measurements. A new numerical algorithm is proposed (1) to capture the actual experimental noise and (2) to substantially speed up previous implementations by exposing more concurrency. Based on the Matrix Over Runtime System Environment numerical library (MORSE), a dynamic scheduler drives all computational stages of the tomographic reconstructor simulation and allows to pipeline and to run tasks out-of-order across different stages on heterogeneous systems, while ensuring data coherency and dependencies.

pTatin3D: High Performance Methods for Long-Term Lithospheric Dynamics

Dave A. May (ETH Zürich), Jed Brown (Argonne National Laboratory), Laetitia Le Pourhiet (Pierre and Marie Curie University)

Simulations of long-term lithospheric deformation involve post-failure analysis of high-contrast brittle materials driven by buoyancy and processes at the free surface. Geodynamic phenomena such as subduction and continental rifting take place over millions year time scales, thus require efficient solution methods. We present pTatin3D, a geodynamics modelling package utilizing the material-point-method for tracking material composition, combined with a multigrid finite-element method to solve heterogeneous, incompressible visco-plastic Stokes problems. Here we analyze the performance and algorithmic tradeoffs of pTatin3D's multigrid preconditioner. Our matrix-free geometric multigrid preconditioner trades flops for memory bandwidth to produce a time-to-solution greater than 2 times faster than the best available methods utilizing stored matrices (plagued by memory bandwidth limitations), exploits local element structure to achieve weak scaling at 30% of FPU peak on Cray XC-30, has improved dynamic range due to smaller memory footprint, and has more consistent timing and better intra-node scalability due to reduced memory-bus and cache pressure.

Compiler Analysis and Optimization

Chair: Milind Kulkarni (Purdue University)

3:30pm-5pm

Room: 393-94-95

Oil and Water Can Mix! An Integration of Polyhedral and AST-based Transformations

Jun Shirako (Rice University), Louis-Noel Pouchet (University of California, Los Angeles), Vivek Sarkar (Rice University)

Optimizing compilers targeting modern multi-core machines require complex program restructuring to expose the proper grain of coarse- and fine-grain parallelism and data locality. The polyhedral compilation model has provided significant advancements in handling of compositions of loop transforma-

tions, exposing multiple levels of parallelism and improving data reuse. However, not all program transformations can be expressed in this model, and some others may actually limit performance because of the excessively complex loop structures generated.

In this paper, we propose an optimization flow that combines polyhedral and syntactic/AST-based transformations, leveraging the strengths and cornering the limitations of each framework. It generates high-performance code containing well-formed loops which can be effectively vectorized, while still exposing sufficient parallelism and data reuse. It combines several transformation stages using either polyhedral or AST-based transformations, delivering performance improvements over single-staged polyhedral compilers.

Award: Best Paper Finalists

Compiler Techniques for Massively Scalable Implicit Task Parallelism

Timothy G. Armstrong (University of Chicago), Justin M. Wozniak, Michael Wilde, Ian T. Foster (University of Chicago and Argonne National Laboratory)

Swift/T is a high-level language for writing concise, deterministic scripts that compose serial or parallel codes implemented in lower-level programming models into large-scale parallel applications. It executes using a data-driven task parallel execution model that is capable of orchestrating millions of concurrently executing asynchronous tasks on homogeneous or heterogeneous resources.

Producing code that efficiently executes at this scale requires sophisticated compiler transformations: poorly optimized code inhibits scaling with excessive synchronization and communication. We present a comprehensive set of compiler techniques for data-driven task parallelism, including novel compiler optimizations and intermediate representations. We report application benchmark studies, including unbalanced tree search and simulated annealing, and demonstrate that our techniques greatly reduce communication overhead and enable extreme scalability, distributing up to 612 million dynamically load balanced tasks per second at scales of up to 262,144 cores without explicit parallelism, synchronization, or load balancing in application code.

MSL: A Synthesis Enabled Language for Distributed Implementations

Zhilei Xu, Shoaib Kamil, Armando Solar-Lezama (Massachusetts Institute of Technology)

This paper presents a new methodology for implementing SPMD distributed memory kernels. In this methodology, the programmer describes a generic implementation strategy for a particular class of kernels by writing a partial, generic implementation and wrapping it as a library with a concise and clean

interface. The programmer then relies on our system to derive a concrete MPI implementation that follows the strategy and matches a given reference implementation. This methodology is made possible by a new set of language constructs that allow programmers to relate the behavior of a sequential reference implementation to a distributed implementation, and a new synthesis algorithm for distributed memory implementations. We demonstrate the methodology by implementing non-trivial kernels from NAS Parallel Benchmarks. Our approach can automatically infer challenging details and produce efficient implementations that perform within 95% of hand-written Fortran code, while promoting reusability and reducing programmer effort by leveraging generative programming and synthesis.

Networks

Chair: Milind Kulkarni (Purdue University)

3:30pm-5pm

Room: 388-89-90

RAHTM: Routing Algorithm Aware Hierarchical Task Mapping

Ahmed H. Abdel-Gawad (Purdue University), Mithuna Thottethodi (Purdue University), Abhinav Bhatele (Lawrence Livermore National Laboratory)

The mapping of MPI processes to compute nodes in a supercomputer has a first-order impact on communication performance. For high-performance computing (HPC) applications with stable communication, rich offline analysis of such communication can improve the communication performance by optimizing the mapping. Unfortunately, current practices for at-scale HPC consider only the communication graph and network topology in solving this problem.

We propose Routing-aware Hierarchical Task Mapping (RAHTM) which leverages routing algorithm awareness to improve task mapping. RAHTM achieves high-quality mappings by combining (1) a divide-and-conquer strategy to achieve scalability, (2) a limited search of mappings, and (3) a linear programming based routing-aware approach to evaluate possible mappings in the search space. RAHTM achieves 20% reduction in communication time and a 9% reduction in overall execution time, for three communication-heavy benchmarks scaled up to 16K processes on a Blue Gene/Q platform.

Maximizing Throughput on a Dragonfly Network

Nikhil Jain (University of Illinois at Urbana-Champaign), Abhinav Bhatele (Lawrence Livermore National Laboratory), Xiang Ni (University of Illinois at Urbana-Champaign), Nicholas J. Wright (Lawrence Berkeley National Laboratory), Laxmikant V. Kale (University of Illinois at Urbana-Champaign)

Interconnection networks are a critical resource for large supercomputers. The dragonfly topology, which provides a low

network diameter and large bisection bandwidth, is being explored as a promising option for building multi-Petaflop/s and Exaflop/s systems. Unlike the extensively studied torus networks, the best choices of message routing and job placement strategies for the dragonfly topology are not well understood. This paper aims at analyzing the behavior of a machine built using a dragonfly network for various routing strategies, job placement policies, and application communication patterns. Our study is based on a novel model that predicts traffic on individual links for direct, indirect, and adaptive routing strategies. We analyze results for individual communication patterns and some common parallel job workloads. The predictions presented in this paper are for a 100+ Petaflop/s prototype machine with 92,160 high-radix routers and 8.8 million cores.

Slim Fly: A Cost Effective Low-Diameter Network Topology

Maciej Besta (ETH Zurich), Torsten Hoefler (ETH Zurich)

We introduce a high-performance cost-effective network topology called Slim Fly that approaches the theoretically optimal network diameter. Slim Fly is based on graphs that approximate the solution to the degree-diameter problem. We analyze Slim Fly and compare it to both traditional and state-of-the-art networks. Our analysis shows that Slim Fly has significant advantages over other topologies in latency, bandwidth, resiliency, cost, and power consumption. Finally, we propose deadlock-free routing schemes and physical layouts for large computing centers as well as a detailed cost and power model. Slim Fly enables constructing cost effective and highly resilient datacenter and HPC networks that offer low latency and high bandwidth under different HPC workloads such as stencil or graph computations.

Award: Best Student Paper Finalists

Parallel Algorithms

Chair: Judith C. Hill (Oak Ridge National Laboratory)

3:30pm-5pm

Room: 391-92

A Computation- and Communication-Optimal Parallel Direct 3-Body Algorithm

Penporn Koanantakool (University of California, Berkeley), Katherine Yelick (Lawrence Berkeley National Laboratory)

Traditional particle simulation methods are used to calculate pairwise potentials, but some problems require 3-body potentials that calculate over triplets of particles. A direct calculation of 3-body interactions involves $O(n^3)$ interactions, but has significant redundant computations that occur in a nested loop formulation. In this paper we explore algorithms for 3-body computations that simultaneously optimize three criteria: computation minimization through symmetries, communication optimality, and load balancing. We present a new 3-body algorithm that is both communication and computa-

tion optimal. Its optional replication factor, c , saves c^3 in latency (number of messages) and c^2 in bandwidth (volume), with bounded load-imbalance. We also consider the k -body case and discuss an algorithm that is optimal if there is a cutoff distance of less than $1/3$ of the domain. The 3-body algorithm demonstrates 99% efficiency on tens of thousands of cores, showing strong scaling properties with order of magnitude speedups over the naïve algorithm.

A Communication-Optimal Framework for Contracting Distributed Tensors

Samyam Rajbhandari, Akshay Nikam, Pai-Wei Lai, Kevin Stock (Ohio State University), Sriram Krishnamoorthy (Pacific Northwest National Laboratory), P. Sadayappan (Ohio State University)

Tensor contractions are extremely compute intensive generalized matrix multiplication operations encountered in many computational science fields, such as quantum chemistry and nuclear physics. Unlike distributed matrix multiplication, which has been extensively studied, limited work has been done in understanding distributed tensor contractions. In this paper, we characterize distributed tensor contraction algorithms on torus networks. We develop a framework with three fundamental communication operators to generate communication-efficient contraction algorithms for arbitrary tensor contractions. We show that for a given amount of memory per processor, our framework is communication optimal for all tensor contractions. We demonstrate performance and scalability of our framework on up to 262,144 cores of BG/Q supercomputer using five tensor contraction examples.

Award: Best Paper Finalists

Fast Parallel Computation of Longest Common Prefixes

Julian Shun (Carnegie Mellon University)

Suffix arrays and the corresponding longest common prefix (LCP) array have wide applications in bioinformatics, information retrieval and data compression. In this work, we propose and theoretically analyze new parallel algorithms for computing the LCP array given the suffix array as input. Most of our algorithms have a work and depth (parallel time) complexity related to the LCP values of the input. We also present a slight variation of Karkkainen and Sanders' skew algorithm that requires linear work and poly-logarithmic depth in the worst case. We present a comprehensive experimental study of our parallel algorithms along with existing parallel and sequential LCP algorithms. On a variety of real-world and artificial strings, we show that on a 40-core shared-memory machine our fastest algorithm is up to 2.3 times faster than the fastest existing parallel algorithm, and up to 21.8 times faster than the fastest sequential LCP algorithm.

Wednesday, November 19

Big Data Analysis

Chair: Kelly Gaither (Texas Advance Computing Center)

10:30am-12pm

Room: 391-92

Fast Iterative Graph Computation:

A Path Centric Approach

Pingpeng Yuan, Wenya Zhang, Changfeng Xie (Huazhong University of Science & Technology), Ling Liu (Georgia Institute of Technology), Hai Jin (Huazhong University of Science & Technology), Kisung Lee (Georgia Institute of Technology)

Large-scale graph processing represents an interesting challenge due to the lack of locality. This paper presents Path-Graph for improving iterative graph computation on graphs with billions of edges. Our system design has three unique features: First, we model a large graph using a collection of tree-based partitions and use a path-centric computation rather than vertex-centric or edge-centric computation. Our parallel computation model significantly improves the memory and disk locality for performing iterative computation algorithms. Second, we design a compact storage that further maximize sequential access and minimize random access on storage media. Third, we implement the path-centric computation model by using a scatter/gather programming model, which parallels the iterative computation at partition tree level and performs sequential updates for vertices in each partition tree. The experimental results show that the path-centric approach outperforms vertex-centric and edge-centric systems on a number of graph algorithms for both in-memory and out-of-core graphs.

Efficient I/O and Storage of Adaptive Resolution Data

Sidharth Kumar, John Edwards (University of Utah), Peer-Timo Bremer (Lawrence Livermore National Laboratory), Aaron Knoll, Cameron Christensen (University of Utah), Venkatram Vishwanath, Philip Carns (Argonne National Laboratory), John A. Schmidt, Valerio Pascucci (University of Utah)

We present an efficient, flexible, adaptive-resolution I/O framework that is suitable for both uniform and Adaptive Mesh Resolution (AMR) simulations. In an AMR setting, current solutions typically represent each resolution level as independent grids which often results in inefficient storage and performance. Our technique coalesces domain data into a unified, multi-resolution representation with fast spatially aggregated I/O. Furthermore, our framework easily extends to importance-driven storage of uniform grids, for example, by storing regions of interest at full resolution and nonessential

regions at lower resolution for visualization or analysis. Our framework, which is an extension of the PIDX framework, achieves state of the art disk usage and I/O performance regardless of resolution of the data, regions of interest, and the number of processes that generated the data. We demonstrate the scalability and efficiency of our framework using the Uintah and S3D large-scale combustion codes on machines Mira and Edison.

An Image-Based Approach to Extreme Scale In Situ Visualization and Analysis

James Ahrens (Los Alamos National Laboratory), Sebastien Jourdain, Patrick O'Leary (Kitware Inc), John Patchett, David Rogers, Mark Petersen (Los Alamos National Laboratory)

Extreme scale scientific simulations are leading a charge to exascale computation, and data analytics runs the risk of being a bottleneck to scientific discovery. Due to power and I/O constraints, we expect in situ visualization and analysis will be a critical component of these workflows. Options for extreme scale data analysis are often presented as a stark contrast: write large files to disk for interactive, exploratory analysis, or perform in situ analysis to save detailed data about phenomena that a scientist knows about in advance. We present a novel framework for a third option – a highly interactive, image-based approach that promotes exploration of simulation results, and is easily accessed through extensions to widely used open source tools. This in situ approach supports interactive exploration of a wide range of results, while still significantly reducing data movement and storage.

High Performance Genomics

Chair: Zhong Jin (Supercomputing Center, Chinese Academy of Sciences)

10:30am-12pm

Room: 388-89-90

Parallel De Bruijn Graph Construction and Traversal for De Novo Genome Assembly

Evangelos Georganas (University of California, Berkeley), Aydin Buluc (Lawrence Berkeley National Laboratory), Jarrod Chapman (Joint Genome Institute), Leonid Olier (Lawrence Berkeley National Laboratory), Daniel Rokhsar (Joint Genome Institute), Katherine Yelick (Lawrence Berkeley National Laboratory)

De novo whole genome assembly reconstructs genomic sequence from short, overlapping, and potentially erroneous fragments called reads. We study optimized parallelization of the most time-consuming phases of Meraculous, a state-of-the-art production assembler. First, we present a new parallel algorithm for k-mer analysis, characterized by intensive communication and I/O requirements, and reduce the memory requirements by 6.93x. Second, we efficiently parallelize de

Bruijn graph construction and traversal, which necessitates a distributed hash table and is a key component of most de novo assemblers. We provide a novel algorithm that leverages one-sided communication capabilities of the Unified Parallel C (UPC) to facilitate the requisite fine-grained parallelism and avoidance of data hazards, while analytically proving its scalability properties. Overall results show unprecedented performance and efficient scaling on up to 15,360 cores of the Cray XC30, on human genome as well as the challenging wheat genome, with performance improvement from days to seconds.

Orion : Scaling Genomic Sequence Matching with Fine-Grained Parallelization

Kanak Mahadik, Somali Chaterji, Bowen Zhou, Milind Kulkarni, Saurabh Bagchi (Purdue University)

Gene sequencing instruments are producing huge volumes of data, straining the capabilities of current database searching algorithms and hindering efforts of researchers analyzing large collections of data to obtain greater insights. In the space of parallel genomic sequence search, most of the popular software packages, like mpiBLAST, use the database segmentation approach, wherein the entire database is sharded and searched on different nodes. However this approach does not scale well with the increasing length of individual query sequences as well as the rapid growth in size of sequence databases. In this paper, we propose a fine-grained parallelism technique called Orion that divides the input query into an adaptive number of fragments and shards the database. Our technique achieves higher parallelism (and hence speedup) and load balancing than database sharding alone, while maintaining 100% accuracy. We show that it is 12.3X faster than mpiBLAST for solving a relevant comparative genomics problem.

Parallel Bayesian Network Structure Learning for Genome-Scale Gene Networks

Sanchit Misra (Intel Corporation), Vasimuddin Md (Indian Institute of Technology Bombay), Kiran Pamnany (Intel Corporation), Sriram P. Chockalingam (Indian Institute of Technology Bombay), Yong Dong, Min Xie (National University of Defense Technology), Maneesha R. Aluru, Srinivas Aluru (Georgia Institute of Technology)

Learning Bayesian networks is NP-hard. Even with recent progress in heuristic and parallel algorithms, modeling capabilities still fall short of the scale of the problems encountered. In this paper, we present a massively parallel method for Bayesian network structure learning, and demonstrate its capability by constructing genome-scale gene networks of the model plant *Arabidopsis thaliana* from over 168.5 million gene expression values. We report strong scaling efficiency of 75% and demon-

strate scaling to 1.57 million cores of the Tianhe-2 supercomputer. Our results constitute three and five orders of magnitude increase over previously published results in the scale of data analyzed and computations performed, respectively. We achieve this through algorithmic innovations, using efficient techniques to distribute work across all compute nodes, all available Intel Xeon processors and Intel Xeon Phi coprocessors on each node, all available threads on each processor and coprocessor, and vectorization techniques to maximize single thread performance.

Award: Best Paper Finalist

MPI

Chair: Ron Brightwell (Sandia National Laboratories)

10:30am-12pm

Room: 393-94-95

Nonblocking Epochs in MPI One-Sided Communication
Judicael A. Zounmevo (Queen's University), Xin Zhao (University of Illinois at Urbana-Champaign), Pavan Balaji (Argonne National Laboratory), William Gropp (University of Illinois at Urbana-Champaign), Ahmad Afsahi (Queen's University)
The synchronization model of the MPI one-sided communication paradigm can lead to serialization and latency propagation. For instance, a process can propagate non RMA communication-related latencies to remote peers waiting in their respective epoch-closing routines in matching epochs. In this work, we discuss six latency issues that were documented for MPI-2.0 and show how they evolved in MPI-3.0. Then, we propose entirely non-blocking RMA synchronizations that allow processes to avoid waiting even in epoch-closing routines. The proposal provides contention avoidance in communication patterns that require back-to-back RMA epochs. It also fixes the aforementioned latency propagation issues. It finally allows the MPI progress engine to orchestrate aggressive scheduling to cut down the overall completion time of sets of epochs without introducing memory consistency hazards. Our test results show noticeable performance improvements for a Lower-Upper matrix decomposition as well as an application pattern that performs massive atomic updates.

Award: Best Paper Finalists

Enabling Efficient Multithreaded MPI Communication through a Library-Based Implementation of MPI Endpoints

Srinivas Sridharan, James Dinan, Dhiraj Kalamkar (Intel Corporation)

Modern high-speed interconnection networks are designed with capabilities to support communication from multiple processor cores. The MPI endpoints extension has been proposed to ease process and thread count tradeoffs by enabling multithreaded MPI applications to efficiently drive independent

network communication. In this work, we present the first implementation of the MPI endpoints interface and demonstrate the first applications running on this new interface. We use a novel library-based design that can be layered on top of any existing, production MPI implementation. Our approach uses proxy processes to isolate threads in an MPI job, eliminating threading overheads within the MPI library and allowing threads to achieve process-like communication performance. Performance results for the Lattice QCD Dslash kernel indicates that endpoints provides up to 2.9x improvement in communication performance and 1.87x overall performance improvement over a highly optimized hybrid MPI+OpenMP baseline on 128 processors.

MC-Checker: Detecting Memory Consistency Errors in MPI One-Sided Applications

Zhezhe Chen (Twitter Inc.), James Dinan (Intel Corporation), Zhen Tang (Institute of Software, Chinese Academy of Sciences), Pavan Balaji (Argonne National Laboratory), Hua Zhong, Jun Wei, Tao Huang (Institute of Software, Chinese Academy of Sciences), Feng Qin (Ohio State University)

One-sided communication decouples data movement and synchronization by providing support for asynchronous reads and updates of distributed shared data. While such interfaces can be extremely efficient, they also impose challenges in properly performing asynchronous accesses to shared data.

This paper presents MC-Checker, a new tool that detects memory consistency errors in MPI one-sided applications. MC-Checker first performs online instrumentation and captures relevant dynamic events, such as one-sided communications and load/store operations. MC-Checker then performs analysis to detect memory consistency errors. When found, errors are reported along with useful diagnostic information. Experiments indicate that MC-Checker is effective at detecting and diagnosing memory consistency bugs in MPI one-sided applications, with low overhead, ranging from 24.6% to 71.1%, with an average of 45.2%.

Cloud Computing I

Chair: Rizos Sakellariou (University of Manchester)

1:30pm-3pm

Room: 391-92

Scheduling Multi-Tenant Cloud Workloads on Accelerator-Based Systems

Dipanjana Sengupta, Anshuman Goswami, Karsten Schwan, Krishna Pallavi (Georgia Institute of Technology)

Accelerator-based systems are making rapid inroads into becoming platforms of choice for high end cloud services. There is a need therefore, to move from the current model in which high performance applications explicitly and programmatically select the GPU devices on which to run, to a dynamic model

where GPUs are treated as first class schedulable entities. The Strings scheduler realizes this vision by decomposing the GPU scheduling problem into a combination of load balancing and per-device scheduling. (i) Device-level scheduling efficiently uses all of a GPU's hardware resources, including its computational and data movement engines, and (ii) load balancing goes beyond obtaining high throughput, to ensure fairness through prioritizing GPU requests that have attained least service. With its methods, Strings achieves system throughput and fairness improvements of up to 8.70x and 13%, respectively, compared to the CUDA runtime

Scaling MapReduce Vertically and Horizontally

Ismail El-Helw, Rutger Hofman, Henri Bal (VU University Amsterdam)

Glasswing is a MapReduce framework that uses OpenCL to exploit multi-core CPUs and accelerators. However, compute device capabilities vary significantly and require targeted optimization. Similarly, availability of memory, storage and interconnects impacts job performance. In this paper, we present and analyze how MapReduce applications can improve their horizontal and vertical scalability using a well controlled mixture of coarse- and fine-grained parallelism. We discuss the Glasswing pipeline and its ability to overlap computation, communication, memory transfers and disk access. We show how Glasswing adapts to the distinct capabilities of a variety of compute devices by employing fine-grained parallelism. We experimentally evaluated the performance of five applications and show that Glasswing outperforms Hadoop on a 64-node multi-core CPU cluster by factors between 1.2 and 4, and factors from 20 to 30 on a 23-node GPU cluster. Similarly, we show that Glasswing is at least 1.5x faster than GPMR on the GPU cluster.

Award: Best Student Paper Finalists

The DRIHM Project: A Flexible Approach to Integrate HPC, Grid and Cloud Resources for Hydro-Meteorological Research

Daniele D'Agostino, Andrea Clematis (National Research Council of Italy, Institute of Applied Mathematics and Information Technology); Dieter Kranzlmüller, Michael Schiffrers, Nils Gentschen Felde, Christian Straube (Ludwig-Maximilians-Universität München); Olivier Caumont (CNRM-GAME (CNRS, Météo-France)), Evelynne Richard (CNRS - University of Toulouse III), Luis Garrote (Universidad Politécnica de Madrid), Quillon Harpham (HR Wallingford Ltd.), H.R.A. Jagers (Deltares); Vladimir Dimitrijević, Ljiljana Dekić (Republic Hydrometeorological Service of Serbia); Elisabetta Fiori, Fabio Delogu, Antonio Parodi (CIMA Research Foundation)

The distributed research infrastructure for hydro-meteorology (DRIHM) project focuses on the development of an e-Science infrastructure to provide end-to-end hydro-meteorological research (HMR) services (models, data, and post-processing

tools) by exploiting HPC, Grid and Cloud facilities. In particular, the DRIHM infrastructure supports the execution and analysis of high-resolution simulations through the definition of workflows composed by heterogeneous HMR models in a scalable and interoperable way, while hiding all the low level complexities. This contribution gives insights into best practices adopted to satisfy the requirements of an emerging multidisciplinary scientific community composed of earth and atmospheric scientists. To this end, DRIHM supplies innovative services leveraging high performance and distributed computing resources. Hydro-meteorological requirements shape this IT infrastructure through an iterative “learning-by-doing” approach that permits tight interactions between the application community and computer scientists, leading to the development of a flexible, extensible, and interoperable framework.

Graph Algorithms

Chair: Felix Wolf (*German Research School for Simulation Sciences*)

1:30pm-3pm

Room: 388-89-90

Faster Parallel Traversal of Scale Free Graphs at Extreme Scale with Vertex Delegates

Roger Pearce, Maya Gokhale (Lawrence Livermore National Laboratory), Nancy M. Amato (Texas A&M University)

At extreme scale, irregularities in the structure of scale-free graphs such as social network graphs limit our ability to analyze these important and growing datasets. A key challenge is the presence of high-degree vertices that leads to parallel workload and storage imbalances. We present techniques to distribute storage, computation, and communication of hubs for extreme scale graphs in distributed memory supercomputers. To balance the hub processing workload, we distribute hub data structures and related computation among a set of delegates. The delegates coordinate using highly optimized asynchronous broadcast and reduction operations. We demonstrate scalability of our new algorithmic technique using Breadth-First Search (BFS), Single Source Shortest Path (SSSP), K-Core Decomposition, and Page-Rank on synthetically generated scale-free graphs. Our results show excellent scalability on large scale-free graphs up to 131K cores of the IBM BG/P, and outperform the best known Graph500 performance on BG/P Intrepid by 15%.

Pardicle: Parallel Approximate Density-Based Clustering

Md. Mostofa Ali Patwary, Nadathur Satish, Narayanan Sundaram (Intel Corporation), Fredrik Manne (University of Bergen, Norway), Salman Habib (Argonne National Laboratory), Pradeep Dubey (Intel Corporation)

DBSCAN is a widely used isodensity-based clustering algorithm for particle-data well-known for its ability to isolate arbitrarily-shaped clusters and to filter noise-data. The algorithm is super-linear ($O(n \log n)$) and computationally expensive for large datasets. Given the need for speed, we propose an approximate DBSCAN algorithm using density-based-sampling, which performs equally well in quality compared to exact algorithms, but is more than an order-of-magnitude faster. Our experiments on astrophysics and synthetic massive-datasets (8.5B numbers) shows that our approximate algorithm is up to 56x faster than exact algorithms with almost identical quality ($\text{Omega-Index} \geq 0.99$). We develop a new parallel DBSCAN algorithm, which uses dynamic-partitioning to improve load-balancing and locality. We demonstrate near-linear speedup on shared memory (15x on 16-core Intel® Xeon® E5-2680 systems and 59x on Intel® Xeon Phi™ with 2x performance improvement over Xeon) and distributed memory (3917x using 4096 Xeon cores) computers. Additionally, existing exact algorithms can achieve up to 3.4 times speedup using dynamic partitioning.

Scalable and High Performance Betweenness Centrality on the GPU

Adam T. McLaughlin, David A. Bader (Georgia Institute of Technology)

Graphs that model social networks, numerical simulations, and the structure of the Internet are enormous and cannot be manually inspected. A popular metric used to analyze these networks is betweenness centrality, which has applications in community detection, power grid contingency analysis, and the study of the human brain. However, these analyses come with a high computational cost that prevents the examination of large graphs of interest.

Prior GPU implementations suffer from large local data structures and inefficient graph traversals that limit scalability and performance. Here we present several hybrid GPU implementations, providing good performance on graphs of arbitrary structure rather than just scale-free graphs as was done previously. We achieve up to 13x speedup on high-diameter graphs and an average of 2.71x speedup overall over the best existing GPU algorithm. We observe near linear speedup and performance exceeding tens of GTEPS when running betweenness centrality on 192 GPUs.

Award: Best Student Paper Finalists

Hardware Vulnerability and Recovery

Chair: Alison Kennedy (Edinburgh Parallel Computing Center)

1:30pm-3pm

Room: 393-94-95

Understanding Soft Error Resiliency of BlueGene/Q Compute Chip through Hardware Proton Irradiation and Software Fault Injection

Chen-Yong Cher, Meeta S. Gupta, Pradip Bose, K. Paul Muller (IBM Systems and Technology Group)

Soft Error Resiliency is a major concern for Petascale high performance computing (HPC) systems. Blue Gene/Q (BG/Q) is the third generation of IBM's massively parallel, energy efficient Blue Gene series of supercomputers. The principal goal of this work is to understand the interaction between BlueGene/Q's hardware resiliency features and high-performance applications through proton irradiation of a real chip, and software resiliency inherent in these applications through application-level fault injection (AFI) experiments. From the proton irradiation experiments we derived that the mean time between correctable errors at sea level of the SRAM-based register files and Level-1 caches for a system similar to the scale of Sequoia system. From the AFI experiments, we characterized relative vulnerability among the applications in both general purpose and floating point register files. We categorized and quantified the failure outcomes, and discovered characteristics in the applications that may lead to many opportunities for improvement of resilience.

Fail-in-Place Network Design: Interaction between Topology, Routing Algorithm and Failures

Jens Domke (Tokyo Institute of Technology), Torsten Hoefler (ETH Zürich), Satoshi Matsuoka (Tokyo Institute of Technology)

The growing system size of high performance computers results in a steady decrease of the mean time between failures. Exchanging network components often requires whole system downtime which increases the cost of failures. In this work, we study a fail-in-place strategy where broken network elements remain untouched. We show that a fail-in-place strategy is feasible for today's networks and the degradation is manageable, and provide guidelines for the design. Our network failure simulation toolchain allows system designers to extrapolate the performance degradation based on expected failure rates, and it can be used to evaluate the current state of a system. In a case study of real-world HPC systems, we will analyze the performance degradation throughout the system's lifetime under the assumption that faulty network components are not repaired, which results in a recommendation to change the used routing algorithm to improve the network performance as well as the fail-in-place characteristic.

Correctness Field Testing of Production and Decommissioned High Performance Computing Platforms at Los Alamos National Laboratory

Sarah E. Michalak, William N. Rust (Los Alamos National Laboratory), John T. Daly (Laboratory for Physical Sciences), Andrew J. DuBois, David H. DuBois (Los Alamos National Laboratory)

Silent Data Corruption (SDC) can threaten the integrity of scientific calculations performed on high performance computing (HPC) platforms and other systems. To characterize this issue, correctness field testing of HPC platforms at Los Alamos National Laboratory was performed. This work presents results for 12 platforms, including over 1,000 node-years of computation performed on over 8,750 compute nodes and over 260 PB of data transfers involving nearly 6,000 compute nodes, and relevant lessons learned. Incorrect results characteristic of transient errors and of intermittent errors were observed. These results are a key underpinning to resilience efforts as they provide signatures of incorrect results observed under field conditions. Five incorrect results consistent with a transient error mechanism were observed, suggesting that the effects of transient errors could be mitigated. However, the observed numbers of incorrect results consistent with an intermittent error mechanism suggest that intermittent errors could substantially affect computational correctness.

I/O and Dynamic Optimization

Chair: John West (DOD HPC Modernization Program)

3:30pm-5pm

Room: 391-92

Omnisc'IO: A Grammar-Based Approach to Spatial and Temporal I/O Patterns Prediction

Matthieu Dorier, Shadi Ibrahim, Gabriel Antoniu (Inria, Rennes Bretagne Atlantique Research Centre), Rob Ross (Argonne National Laboratory)

The increasing gap between the computation performance of post-petascale machines and the performance of their I/O subsystem has motivated many I/O optimizations including prefetching, caching, and scheduling techniques. In order to further improve these techniques, modeling and predicting spatial and temporal I/O patterns of HPC applications as they run has become crucial. In this paper we present Omnisc'IO, an approach that builds a grammar-based model of the I/O behavior of HPC applications and uses it to predict when future I/O operations will occur, and where and how much data will be accessed. Omnisc'IO is transparently integrated into the POSIX and MPI I/O stacks and does not require any modification in applications or higher-level I/O libraries. It works without any prior knowledge of the application and converges to accurate predictions within a couple of iterations only. Its implementation is efficient in both computation time and memory footprint.

Two-Choice Randomized Dynamic I/O Scheduler for Object Storage Systems

Dong Dai, Yong Chen (Texas Tech University), Dries Kimpe, Rob Ross (Argonne National Laboratory)

Object storage is considered a promising solution for the next-generation (exascale) high performance computing platform due to its flexible and high performance object interface. However, delivering high burst-write throughput is still a critical challenge. Although deploying more storage servers can potentially provide higher throughput, it can be ineffective as the burst-write throughput is limited by a small number of possible stragglers. In this paper, we propose a two-choice randomized dynamic I/O scheduler to avoid stragglers and hence achieve high throughput. The contributions include:

1) we propose a two-choice randomized dynamic I/O scheduler with collaborative probe and preassign strategies; 2) we design and implement a redirect table and metadata maintainer to address the metadata management challenge introduced by dynamic I/O scheduling; and 3) we evaluate the proposed scheduler with both simulation and experimental tests in HPC cluster. The evaluation results confirmed the scalability and performance benefits of the proposed I/O scheduler.

Parallel Programming with Migratable Objects: Charm++ in Practice

Bilge Acun, Abhishek Gupta, Nikhil Jain, Akhil Langer, Harshitha Menon, Eric Mikida, Xiang Ni, Michael Robson, Yanhua Sun, Ehsan Totoni, Lukasz Wesolowski, Laxmikant Kale (University of Illinois at Urbana-Champaign)

The advent of petascale computing has introduced new challenges (e.g., heterogeneity, system failure) for programming scalable parallel applications. Increased complexity and dynamism in science and engineering applications of today have further exacerbated the situation. Addressing these challenges requires more emphasis on concepts that were previously of secondary importance, including migratability, adaptivity, and runtime system introspection. In this paper, we leverage our experience with these concepts to demonstrate their applicability and efficacy for real world applications. Using the CHARM++ parallel programming framework, we present details on how these concepts can lead to development of applications that scale irrespective of the rough landscape of supercomputing technology. Empirical evaluation presented in this paper spans many mini-applications and real applications executed on modern supercomputers including Blue Gene/Q, Cray XE6, and Stampede.

Quantum Simulations in Materials and Chemistry

Chair: Jack Wells (Oak Ridge National Laboratory)

3:30pm-5pm

Room: 388-89-90

Metascaleable Quantum Molecular Dynamics Simulations of Hydrogen-on-Demand

Ken-ichi Nomura, Rajiv K. Kalia, Aiichiro Nakano, Priya Vashishta (University of Southern California); Kohei Shimamura, Fuyuki Shimojo (Kumamoto University); Manaschai Kunaseth (National Nanotechnology Center); Paul C. Messina, Nichols A. Romero (Argonne National Laboratory)

We enabled an unprecedented scale of quantum molecular dynamics simulations through algorithmic innovations. A new lean divide-and-conquer density functional theory algorithm significantly reduces the prefactor of the $O(N)$ computational cost based on complexity and error analyses. A globally scalable and locally fast solver hybridizes a global real-space multigrid with local plane-wave bases. The resulting weak-scaling parallel efficiency was 0.984 on 786,432 IBM Blue Gene/Q cores for a 50.3 million-atom (39.8 trillion degrees-of-freedom) system. The time-to-solution was 60-times less than the previous state-of-the-art, owing to enhanced strong scaling by hierarchical band-space-domain decomposition and high floating-point performance (50.5% of the peak). Production simulation involving 16,661 atoms for 21,140 time steps (or 129,208 self-consistent-field iterations) revealed a novel nanostructural design for on-demand hydrogen production from water, advancing renewable energy technologies. This metascaleable (or “design once, scale on new architectures”) algorithm is used for broader applications within a recently proposed divide-conquer-recombine paradigm.

Efficient Implementation of Many-Body Quantum Chemical Methods on the Intel Xeon Phi Coprocessor

Edoardo Apra (Pacific Northwest National Laboratory), Michael Klemm (Intel Corporation), Karol Kowalski (Pacific Northwest National Laboratory)

This paper presents the implementation and performance of the highly accurate CCSD(T) quantum chemistry method on the Intel Xeon Phi coprocessor within the context of the NWChem computational chemistry package. The widespread use of highly correlated methods in electronic structure calculations is contingent upon the interplay between advances in theory and the possibility of utilizing the ever-growing computer power of emerging heterogeneous architectures. We discuss the design decisions of our implementation as well as the optimizations applied to the compute kernels and data transfers between host and coprocessor. We show the feasibility of adopting the Intel Many Integrated Core Architecture and the Intel Xeon Phi coprocessor for developing efficient

computational chemistry modeling tools. Remarkable scalability is demonstrated by benchmarks. Our solution scales up to a total of 62,560 cores with the concurrent utilization of Intel Xeon processors and Intel Xeon Phi coprocessors.

Optimized Scheduling Strategies for Hybrid Density Functional Theory Electronic Structure Calculations

William DF Dawson, Francois Gygi (University of California, Davis)

Hybrid Density Functional Theory (DFT) has recently gained popularity as an accurate model of electronic interactions in chemistry and materials science applications. The most computationally expensive part of hybrid DFT simulations is the calculation of exchange integrals between pairs of electrons. We present strategies to achieve improved load balancing and scalability for the parallel computation of these integrals. First, we develop a cost model for the calculation, and utilize random search algorithms to optimize the data distribution and calculation schedule. Second, we further improve performance using partial data-replication to increase data availability across cores. We demonstrate these improvements using an implementation in the Qbox Density Functional Theory code on the Mira BlueGene/Q computer at Argonne National Laboratory. We perform calculations in the range of 8k to 128k cores on two representative simulation samples from materials science and chemistry applications: liquid water and a metal-water interface.

Resilience

Chair: Ananta Tiwari (PMaC Lab, San Diego Supercomputer Center)

3:30pm-5pm

Room: 393-94-95

Quantitatively Modeling Application Resiliency with the Data Vulnerability Factor

Li Yu (Illinois Institute of Technology), Dong Li, Sparsh Mittal, Jeffery S. Vetter (Oak Ridge National Laboratory)

Strategies to improve the visible resilience of applications require the ability to distinguish vulnerability difference across application components and selectively apply protection. Hence, quantitatively modeling application vulnerability, as a method to capture vulnerability variance within the application, is critical to evaluate and improve system resilience. The tradition methods cannot effectively quantify vulnerability, because they lack a holist view to examine system resilience, and come with prohibitive evaluation costs. In this paper, we introduce a data-driven methodology to analyze application vulnerability based on a novel resilience metric, the data vulnerability factor (DVF). DVF integrates both application and specific hardware into the resilience analysis. To calculate DVF, we extend a performance modeling language to provide a fast modeling solution. Furthermore, we measure six representa-

tive computational kernels; we demonstrate the values of DVF by quantifying the impact of algorithm optimization on vulnerability and by quantifying the effectiveness of a hardware protection mechanism.

Award: Best Student Paper Finalists

A System Software Approach to Proactive Memory-Error Avoidance

Carlos H. A. Costa Yoonho Park, Bryan S. Rosenburg, Chen-Yong, Kyung Dong Ryu (IBM T. J. Watson Research Center)

Today's HPC systems use two mechanisms to address main-memory errors. Error-correcting codes make correctable errors transparent to software, while checkpoint/restart (CR) enables recovery from uncorrectable errors. Unfortunately, CR overhead will be enormous at exascale due to the high failure rate of memory. We propose a new OS-based approach that proactively avoids memory errors using prediction. This scheme exposes correctable error information to the OS, which migrates pages and offlines unhealthy memory to avoid application crashes. We analyze memory error patterns in extensive logs from a BG/P system and show how correctable error patterns can be used to identify memory likely to fail. We implement a proactive memory management system on BG/Q by extending the firmware and Linux. We evaluate our approach with a realistic workload and compare our overhead against CR. We show improved resilience with negligible performance overhead for applications.

Fault-Tolerant Dynamic Task Graph Scheduling

Mehmet Can Kurt (Ohio State University), Sriram Krishnamoorthy (Pacific Northwest National Laboratory), Kunal Agrawal (Washington University in St. Louis), Gagan Agrawal (Ohio State University)

In this paper, we present an approach to fault-tolerant execution of dynamic task graphs scheduled using work-stealing. In particular, we focus on selective and localized recovery of tasks in the presence of soft faults. We present an application programming interface that elicits the basic task graph structure in terms of successor and predecessor relationships. The work-stealing based task graph scheduling algorithm is then augmented to enable recovery when the data and meta-data associated with a task get corrupted. We use this redundancy, and the information on the task graph provided by the API, to selectively recovery from faults with low space and time overheads. We show that the fault-tolerant design retains the essential properties of the underlying work stealing-based task scheduling algorithm. We also show that the fault tolerant execution is asymptotically optimal when the re-execution of tasks is taken into account. Experimental evaluation demonstrates the effectiveness of the approach.

Award: Best Student Paper Finalists

Thursday, November 20

Machine Learning and Data Analytics

Chair: Hank Childs (University of Oregon and Lawrence Berkeley National Laboratory)

10:30am-12pm

Room: 388-89-90

NUMARCK: Machine Learning Algorithm for Resiliency and Checkpointing

Zhengzhang Chen, Seung Woo Son, William Hendrix, Ankit Agrawal, Wei-keng Liao, Alok Choudhary (Northwestern University)

Data checkpointing is an important fault tolerance technique in HPC systems. This paper exploits the fact that in many scientific applications, relative change in data values from one simulation iteration to the next are not very significantly different from each other. Thus, capturing the distribution of relative changes in data instead of storing data itself allows us to incorporate the temporal dimension of the data, and learn evolving distribution of the changes. We show that an order of magnitude data reduction becomes achievable with a user-defined and guaranteed error bounds for each data point. We propose NUMARCK, NU Machine learning Algorithm for Resiliency and Checkpointing, that makes use of the emerging distributions of data changes between consecutive simulation iterations, and encodes them into an indexing space that can be concisely represented. We evaluate NUMARCK using two production scientific simulations, FLASH and CMIP5, and demonstrate a superior performance.

Parallel Deep Neural Network Training for Big Data on Blue Gene/Q

I-Hsin Chung, Tara Sainath, Bhuvana Ramabhadran, Michael Picheny, John Gunnels, Vernon Austel, Upendra Chauhari, Brain Kingsbury (IBM)

Deep Neural Networks (DNNs) have recently been shown to significantly outperform existing machine learning techniques in several pattern recognition tasks. The biggest drawback to DNNs is the enormous training time - often 10x slower than conventional technologies. While training time can be mitigated by parallel computing algorithms and architectures, these algorithms often suffer from the cost of inter-processor communication bottlenecks. In this paper, we describe how to enable parallel DNN training on the IBM Blue Gene/Q (BG/Q) computer system using the data-parallel Hessian-free 2nd-order optimization algorithm. BG/Q, with its excellent inter-processor communication characteristics, is an ideal match for the HF algorithm. The paper discusses how issues regarding programming model and data-dependent imbalances are addressed. Results on large-scale speech tasks show that the

performance on BG/Q scales linearly up to 4096 processes with no loss in accuracy, allowing us to train DNNs with millions of training examples in a few hours.

FAST: Near Real-time Searchable Data Analytics for the Cloud

Yu Hua (Huazhong University of Science and Technology), Hong Jiang (University of Nebraska-Lincoln), Dan Feng (Huazhong University of Science and Technology)

With the explosive growth in data volume and complexity and the increasing need for highly efficient searchable data analytics, existing cloud storage systems have largely failed to offer an adequate capability for real-time data analytics. To address this problem, we propose a near-real-time and cost-effective searchable data analytics methodology, called FAST. The idea behind FAST is to explore and exploit the semantic correlation within and among datasets via correlation-aware hashing and manageable flat-structured addressing to significantly reduce the processing latency, while incurring acceptably small loss of data-search accuracy. FAST supports several types of data analytics, which can be implemented in existing searchable storage systems. We conduct a real-world use case in which children reported missing in an extremely crowded environment are identified in a timely fashion by analyzing 60 million images using FAST. Extensive experimental results demonstrate the efficiency and efficacy of FAST in the performance improvements and energy savings.

Numerical Kernels

Chair: Kirk E. Jordan (IBM Corporation)

10:30am-12pm

Room: 391-92

Efficient Sparse Matrix-Vector Multiplication on GPUs using the CSR Storage Format

Joseph L. Greathouse, Mayank Daga (Advanced Micro Devices, Inc.)

The performance of sparse matrix vector multiplication (SpMV) is important to computational scientists. Compressed sparse row (CSR) is the most frequently used format to store sparse matrices. However, CSR-based SpMV on graphics processing units (GPUs) has poor performance due to irregular memory access patterns, load imbalance, and reduced parallelism. This has led researchers to propose new storage formats. Unfortunately, dynamically transforming CSR into these formats has significant runtime and storage overheads.

We propose a novel algorithm, CSR-Adaptive, which keeps the CSR format intact and maps well to GPUs. Our implementation addresses the aforementioned challenges by (i) efficiently accessing DRAM by streaming data into the local scratchpad memory and (ii) dynamically assigning different numbers

of rows to each parallel GPU compute unit. CSR-adaptive achieves an average speedup of 14.7x over existing CSR-based algorithms and 2.3x over cISpMV cocktail, which uses an assortment of matrix formats.

Fast Sparse Matrix-Vector Multiplication on GPUs for Graph Applications

Arash Ashari, Naser Sedaghati, John Eisenlohr, Srinivasan Parthasarathy, P. Sadayappan (Ohio State University)

Sparse matrix-vector multiplication (SpMV) is a widely used computational kernel. In this paper, we present ACSR, an adaptive SpMV algorithm that uses the standard CSR format but reduces thread divergence by combining rows into groups (bins) which have a similar number of non-zero elements. Further, for rows in bins that span a wide range of nonzero counts, dynamic parallelism is leveraged. A significant benefit of ACSR over other proposed SpMV approaches is that it works directly with the standard CSR format, and thus avoids significant pre-processing overheads. A CUDA implementation of ACSR is shown to outperform SpMV implementations in the NVIDIA CUSP and cuSPARSE libraries on a set of sparse matrices representing power-law graphs. We also demonstrate the use of ACSR for the analysis of dynamic graphs, where the improvement over extant approaches is even higher.

A Study on Balancing Parallelism and Data Locality in Stencil Calculations

Catherine Olschanowsky, Stephen Guzik (Colorado State University), John Loffeld, Jeffrey Hittinger (Lawrence Livermore National Laboratory), Michelle Strout (Colorado State University)

Large-scale, PDE-based scientific applications are commonly parallelized across large compute resources using MPI. However, the compute power of the resource as a whole can only be utilized if each multicore node is fully utilized. Currently, many PDE solver frameworks parallelize over boxes. In the Chombo framework, the box sizes are typically 16^3 , but larger box sizes such as 128^3 would result in less ghost cell overhead. Unfortunately, typical on-node parallel scaling performs quite poorly for these larger box sizes. In this paper, we investigate around 30 different inter-loop optimization strategies and demonstrate the parallel scaling advantages of some of these variants on NUMA multicore nodes. Shifted, fused, and communication-avoiding variants for 128^3 boxes result in close to ideal parallel scaling and come close to matching the performance of 16^3 boxes on three different multicore systems for an exemplar for many Computational Fluid Dynamic (CFD) codes.

Power and Energy Efficiency

Chair: Kirk Cameron (Virginia Polytechnic Institute and State University)

10:30am-12pm

Room: 393-94-95

Maximizing Throughput of Overprovisioned HPC Data Centers Under a Strict Power Budget

Osman Sarood, Akhil Langer, Abhishek Gupta, Laxmikant Kale (University of Illinois at Urbana-Champaign)

Increasing power consumption is one of the biggest challenges towards achieving the next scale of supercomputers. In this paper, we propose an integer programming based online resource manager, PARM, for data centers with a fixed power budget. PARM determines the optimal allocation of resources to the jobs including CPU power and number of nodes. It leverages the hardware capability of constraining power consumption of nodes, in such a way that the throughput of the data center is maximized for a given power budget. PARM is also capable of dynamically shrinking and expanding the number of nodes of malleable running jobs - a characteristic that further improves its throughput. For resource allocation, it uses the essential power characteristics of the jobs determined by the proposed power-aware performance model. Our results show up to 5.2X improvement in throughput when compared to the power-oblivious SLURM resource manager.

Application Centric Energy-Efficiency Study of Distributed Multi-Core and Hybrid CPU-GPU Systems

Ben Cumming, Gilles Fourestey (Swiss National Supercomputing Centre, ETH Zurich), Oliver Fuhrer (Federal Office of Meteorology and Climatology MeteoSwiss), Tobias Gysi (Department of Computer Science, ETH Zurich), Massimiliano Fatica (NVIDIA Corporation), Thomas C. Schulthess (Swiss National Supercomputing Centre, ETH Zurich)

We study the energy used by a production-level regional climate and weather simulation code on a distributed memory system with hybrid CPU-GPU nodes. The code is optimized for both processor architectures, for which we investigate both time and energy to solution. Operational constraints for time to solution can be met with both processor types, although on different numbers of nodes. Energy to solution is a factor 3 lower on a GPU-system, but strong scaling can be pushed to larger node counts on the CPU subsystem, yielding better time to solution. Our data shows that an affine relationship exists between energy and node-hours consumed by the simulation. We use this property to devise a simple and practical methodology for optimizing for energy efficiency that can be applied to other applications, which we demonstrate with the HPCG benchmark. We conclude with a discussion about the relationship to the commonly-used GF/Watt metric.

Scaling the Power Wall: A Path to Exascale

Oreste Villa, Daniel R. Johnson, Mike O'Connor, Evgeny Bolotin, David W. Nellans, Justin Luitjens, Nikolai Sakharnykh, Peng Wang, Paulius Mickevicius, Anthony Scudiero, Stephen W. Keckler, William J. Dally (NVIDIA Corporation)

Modern scientific discovery is driven by an insatiable demand for computing performance. The HPC community is targeting development of supercomputers able to sustain 1 ExaFlops by the year 2020 and power consumption is the primary obstacle to achieving this goal. A combination of architectural improvements, circuit design, and manufacturing technologies must provide over a 20x improvement in energy efficiency. In this paper, we present some of the progress NVIDIA Research is making toward the design of exascale systems by tailoring features to address the scaling challenges of performance and energy efficiency. We evaluate several architectural concepts for a set of HPC applications demonstrating expected energy efficiency improvements resulting from circuit and packaging innovations such as low-voltage SRAM, low-energy signaling, and on-package memory. Finally, we discuss the scaling of these features with respect to future process technologies and provide power and performance projections for our exascale research architecture.

Data Locality and Load Balancing

Chair: Didem Unat (Koç University)

1:30pm-3pm

Room: 388-89-90

Structure Slicing: Extending Logical Regions with Fields

Michael Bauer, Sean Treichler, Elliott Slaughter, Alex Aiken (Stanford University)

Applications on modern supercomputers are increasingly limited by the cost of data movement, but mainstream programming systems have few abstractions for describing the structure of a program's data. Consequently, the burden of managing data movement, placement, and layout currently falls primarily upon the programmer.

To address this problem we previously proposed a data model based on logical regions and described Legion, a programming system incorporating logical regions. In this paper, we present structure slicing, which incorporates fields into the logical region data model. We show that structure slicing enables Legion to automatically infer task parallelism from field non-interference, decouple the specification of data usage from layout, and reduce the overall amount of data moved. We demonstrate that structure slicing enables both strong and weak scaling of three Legion applications including S3D, a production combustion simulation that uses logical regions with thousands of fields.

Optimizing Data Locality for Fork/Join Programs Using Constrained Work Stealing

Jonathan Lifflander (University of Illinois at Urbana-Champaign), Sriram Krishnamoorthy (Pacific Northwest National Laboratory), Laxmikant Kale (University of Illinois at Urbana-Champaign)

We present an approach to improving data locality across different phases of fork/join programs scheduled using work stealing. The approach consists of: (1) user-specified and automated approaches to constructing a steal tree, the schedule of steal operations and (2) constrained work stealing algorithms that constrain the actions of the scheduler to mirror a given steal tree. These are combined to construct work stealing schedules that maximize data locality across computation phases while ensuring load balance within each phase. These algorithms are also used to demonstrate dynamic coarsening, an optimization to improve spatial locality and sequential overheads by combining many finer-grained tasks into coarser tasks while ensuring sufficient concurrency for locality-optimized load balance. Implementation and evaluation in Cilk demonstrate performance improvements of up to 2.5x on 80 cores. We also demonstrate that dynamic coarsening can combine the performance benefits of coarse task specification with the adaptability of finer tasks.

DISC: A Domain-Interaction Based Programming Model With Support for Heterogeneous Execution

Mehmet Can Kurt (Ohio State University), Gagan Agrawal (Ohio State University)

Several emerging trends are pointing to increasing heterogeneity among nodes and/or cores in HPC systems. Existing programming models, especially for distributed memory execution, typically have been designed to facilitate high performance on homogeneous systems. This paper describes a programming model and an associated runtime system we have developed to address the above need. The main concepts in the programming model are that of a domain and interactions between the domain elements. We explain how stencil computations, unstructured grid computations, and molecular dynamics applications can be expressed using these simple concepts. We show how inter-process communication can be handled efficiently at runtime just from the knowledge of domain interaction, for different types of applications. Subsequently, we develop techniques for the runtime system to automatically partition and re-partition the work among heterogeneous processors or nodes.

Optimized Checkpointing

Chair: Patrick Bridges (University of New Mexico)

1:30pm-3pm

Room: 393-94-95

Understanding the Effects of Communication and Coordination on Checkpointing at Scale

Kurt B. Ferreira (Sandia National Laboratories), Scott Levy (University of New Mexico), Patrick M. Widener (Sandia National Laboratories), Dorian C. Arnold (University of New Mexico), Torsten Hoefler (ETH-Zurich)

Fault-tolerance poses a major challenge for future large-scale systems. However, few insights into selection and tuning of these protocols for applications at scale have emerged. In this paper, we use a simulation-based approach to show that local checkpoint activity in resilience mechanisms can significantly affect the performance of key workloads, even when less than 1% of a local node's compute time is allocated to resilience mechanisms (a very generous assumption). Specifically, we show that even though much work on uncoordinated checkpointing has focused on optimizing message log volumes, local checkpointing activity may dominate the overheads of this technique at scale. Our study shows that local checkpoints lead to process delays that can propagate through messaging relations to other processes causing a cascading series of delays. Lastly, we demonstrate how to tune hierarchical uncoordinated checkpointing protocols designed to reduce log volumes to significantly reduce these synchronization overheads at scale.

Exploring Automatic, Online Failure Recovery for Scientific Applications at Extreme Scales

Marc Gamell (Rutgers University), Daniel S. Katz (University of Chicago and Argonne National Laboratory), Hemanth Kolla (Sandia National Laboratories), Jacqueline Chen (Sandia National Laboratories), Scott Klasky (Oak Ridge National Laboratory), Manish Parashar (Rutgers University)

Application resilience is a key challenge that must be addressed in order to realize the exascale vision. Process/node failures, an important class of failures, are typically handled today by terminating the job and restarting it from the last stored checkpoint. This approach is not expected to scale to exascale. In this paper we present Fenix, a framework for enabling recovery from process/node/blade/cabinet failures for MPI-based parallel applications in an online (i.e., without disrupting the job) and transparent manner. Fenix provides mechanisms for transparently capturing failures, re-spawning new processes, fixing failed communicators, restoring application state, and returning execution control back to the application. To enable automatic data recovery, Fenix relies on application-driven, diskless, implicitly-coordinated check-

pointing. Using the S3D combustion simulation running on the Titan Cray-XK7 production system at ORNL, we experimentally demonstrate Fenix's ability to tolerate high failure rates (e.g., more than one per minute) with low overhead while sustaining performance.

Optimization of Multi-Level Checkpoint Model with Uncertain Execution Scales

Sheng Di (French Institute for Research in Computer Science and Automation and Argonne National Laboratory), Leonardo Bautista-Gomez, Franck Cappello (Argonne National Laboratory)

Future extreme-scale systems are expected to experience different types of failures affecting applications with different failure scales, from transient uncorrectable memory errors in processes to massive system outages. In this paper, we propose a multilevel checkpoint model by taking into account uncertain execution scales (different numbers of processes/cores). The contribution is threefold: (1) we provide an in-depth analysis on why it is difficult to derive the optimal checkpoint intervals for different checkpoint levels and optimize the number of cores simultaneously; (2) we devise a novel method that can quickly obtain an optimized solution - the first successful attempt in multilevel checkpoint models with uncertain scales; and (3) we perform both large-scale real experiments and extreme-scale numerical simulation to validate the effectiveness of our design. The experiments confirm that our optimized solution outperforms other state-of-the-art solutions by 4.3-88% on wall-clock length.

Sparse Solvers

Chair: Anne C. Elster (Norwegian University of Science & Technology/University of Texas at Austin)

1:30pm-3pm

Room: 391-92

Parallelization of Reordering Algorithms for Bandwidth and Wavefront Reduction

Konstantinos I. Karantasis (University of Illinois at Urbana-Champaign), Andrew Lenharth, Donald Nguyen (University of Texas at Austin), Maria Garzaran (University of Illinois at Urbana-Champaign), Keshav Pingali (University of Texas at Austin)

Many sparse matrix computations can be speeded up by first suitably reordering the sparse matrix. Reorderings, originally developed to assist direct methods, have recently become popular for improving cache locality of parallel iterative solvers. When sparse matrix-vector product (SpMV), the key kernel in iterative solvers, runs, reorderings targeting bandwidth and wavefront reduction can improve locality of reference to the compressed sparse storage vectors.

In this paper, we present the first parallel implementations of two widely used reordering algorithms: Reverse Cuthill-McKee (RCM) and Sloan. On 16 cores on Stampede supercomputer, our results show average improvement of 5.56X compared to the state-of-the-art sequential RCM implementation in the HSL library. Sloan is significantly more constrained than RCM; our parallel implementation achieves 2.88X average improvement compared to HSL-Sloan. Considering end-to-end times, using our reordering, SpMV obtains performance improvement of 1.5X compared to natural ordering, and 2X compared to HSL RCM for 100 SpMV iterations.

Domain Decomposition Preconditioners for Communication-Avoiding Krylov Methods on a Hybrid CPU/GPU Cluster

Ichitaro Yamazaki (University of Tennessee, Knoxville), Sivasankaran Rajamanickam, Erik G. Boman, Mark Hoemmen, Michael A. Heroux (Sandia National Laboratories), Stanimire Tomov (University of Tennessee, Knoxville)

Krylov methods are the most widely-used iterative methods for solving large-scale linear systems of equations. Recently, researchers have demonstrated the potential of techniques to avoid communication and improve the performance of the Krylov methods on modern computers, where communication is becoming increasingly expensive. In practice, it is essential to precondition the Krylov method to accelerate its solution convergence. However, we still lack an effective preconditioner to work seamlessly with communication avoiding Krylov methods. We address this crucial gap by presenting a simple communication-avoiding preconditioner. Our preconditioner is based on domain decomposition and does not incur any additional communication than a communication avoiding Krylov method. We illustrate the importance and challenge of developing such preconditioners, and provide the foundation for a family of preconditioners that will be suitable for communication-avoiding methods. Our experimental results with GMRES on distributed GPUs demonstrate the potential of the proposed technique.

Efficient Shared-Memory Implementation of High-Performance Conjugate Gradient Benchmark and Its Application to Unstructured Matrices

Jongsoo Park, Mikhail Smelyanskiy, Karthikeyan Vaidyanathan, Alexander Heinecke, Dhiraj D. Kalamkar (Intel Corporation), Xing Liu (Georgia Institute of Technology), Mostofa Ali Patwary (Intel Corporation), Yutong Lu (National University of Defense Technology), Pradeep Dubey (Intel Corporation)

High performance sparse linear solvers, the backbone of modern HPC, face many challenges on upcoming extreme-scale architectures. The High Performance Linpack (HPL), widely recognized benchmark for ranking such system, does not

represent challenges inherent to these solvers. To address this shortcoming, a new sparse high performance conjugate gradient benchmark (HPCG) has been recently proposed. This is the first paper which analyzes and optimizes HPCG on two modern multi- and many-core IA-based architectures: Xeon and Xeon Phi. We explore number of algorithmic and performance optimizations. By taking advantage of salient architectural features of these two architectures, our implementation sustains 75% and 67% of their achievable bandwidth, respectively. We further show our optimizations generally apply to a wide range of matrices, on which we achieve 72% and 65% of achievable bandwidth. Lastly, we study multi-node scalability of HPCG and the tradeoff between number of parallel domains, convergence and single-node parallel performance.

Cloud Computing II

Chair: Jennifer M. Schopf (Indiana University)

3:30pm-5pm

Room: 388-89-90

FlexSlot: Moving Hadoop into the Cloud with Flexible Slot Management

Yanfei Guo, Jia Rao (University of Colorado at Colorado Springs), Changjun Jiang (Tongji University), Xiaobo Zhou (University of Colorado at Colorado Springs)

Load imbalance is a major source of overhead in Hadoop where the uneven distribution of input data among tasks can significantly delay job completion. Running Hadoop in a private cloud opens up opportunities for mitigating data skew with elastic resource allocation, where stragglers are expedited with more resources, yet introduces problems that often cancel out the performance gain: (1) performance interference from co-running jobs may create new stragglers; (2) there exists a semantic gap between Hadoop task management and resource pool based virtual cluster management preventing efficient resource usage.

We present FlexSlot, a user-transparent task slot management scheme that automatically identifies map stragglers and resizes their slots accordingly to accelerate task execution. FlexSlot adaptively changes the number of slots on each virtual node to promote efficient usage of resource pool. Experimental results with representative benchmarks show that FlexSlot effectively reduces job completion time by 46% and achieves better resource utilization.

Reciprocal Resource Fairness: Towards Cooperative Multiple-Resource Fair Sharing in IaaS Clouds

Haikun Liu (Nanyang Technological University), Bingsheng He (Nanyang Technological University)

Resource sharing in virtualized environments has been demonstrated significant benefits to improve application performance and resource/energy efficiency. However, resource sharing, especially for multiple types, poses several severe and challenging problems in pay-as-you-use cloud environments, such as sharing incentive, economic fairness and free-riding. To address these problems, we propose Reciprocal Resource Fairness (RRF), a novel resource allocation mechanism to enable fair sharing multiple types of resources among multiple tenants in new-generation clouds. RRF implements two complementary and hierarchical mechanisms for resource sharing: inter-tenant resource trading and intra-tenant weight adjustment. We show that RRF satisfies several highly desirable properties to ensure fairness. Experimental results show RRF is promising for both cloud providers and tenants. Compared to the current cloud models, RRF improves VM density and cloud providers' revenue by 2.2X. For tenants, RRF improves application performance by 45% and guarantees 95% economic fairness among multiple tenants.

Finding Constant From Change: Revisiting Network Performance Aware Optimizations on IaaS Clouds

Yifan Gong (NEWRI, Interdisciplinary Graduate School, Nanyang Technological University), Bingsheng He (School of Computer Engineering, Nanyang Technological University), Dan Li (Tsinghua University)

Network performance aware optimizations have long been an effective approach to optimize distributed applications on traditional network environments. But the assumptions of network topology or direct use of several measurements of pair-wise network performance for optimizations are no longer valid on IaaS clouds. Virtualization hides network topology from users, and direct use of network performance measurements may not represent long-term performance. To enable existing network performance aware optimizations on IaaS clouds, we propose to decouple constant component from dynamic network performance while minimizing the difference by a mathematical method called RPCA. We use the constant component to guide network performance aware optimizations and demonstrate the efficiency of our approach by adopting network aware optimizations for collective communications of MPI and generic topology mapping as well as two real-world applications, N-body and conjugate gradient (CG). Our experiments on Amazon EC2 and simulations demonstrate significant performance improvement on guiding the optimizations.

Large-Scale Visualization

Chair: Brent Welch (Google)

3:30pm-5pm

Room: 391-92

High Performance Computation of Distributed-Memory Parallel 3D Voronoi and Delaunay Tessellation

Tom Peterka (Argonne National Laboratory), Dmitriy Morozov (Lawrence Berkeley National Laboratory), Carolyn Phillips (Argonne National Laboratory)

Computing a Voronoi or Delaunay tessellation from a set of points is a core part of the analysis of many simulated and measured datasets: N-body simulations, molecular dynamics codes, and LIDAR point clouds are just a few examples. Such computational geometry methods are common in data analysis and visualization, but as the scale of simulations and observations surpasses billions of particles, the existing serial and shared-memory algorithms no longer suffice. A distributed-memory scalable parallel algorithm is the only feasible approach. The primary contribution of this paper is a new parallel Delaunay and Voronoi tessellation algorithm that automatically determines which neighbor points need to be exchanged among the subdomains of a spatial decomposition. Other contributions include the addition of periodic and wall boundary conditions, comparison of parallelization based on two popular serial libraries, and application to numerous science datasets.

Scalable Computation of Stream Surfaces on Large Scale Vector Fields

Kewei Lu (Ohio State University), Han-Wei Shen (Ohio State University), Tom Peterka (Argonne National Laboratory)

Stream surfaces and streamlines are two popular methods for visualizing three-dimensional flow fields. While several parallel streamline computation algorithms exist, relatively little research has been done to parallelize stream surface generation. This is because load-balanced parallel stream surface computation is nontrivial, due to the strong dependency in computing the positions of the particles forming the stream surface front. In this paper, we present a new algorithm that computes stream surfaces efficiently. In our algorithm, seeding curves are divided into segments, which are then assigned to the processes. Each process is responsible for integrating the segments assigned to it. To ensure a balanced computational workload, work stealing and dynamic refinement of seeding curve segments are employed to improve the overall performance. We demonstrate the effectiveness of our parallel stream surface algorithm using several large scale flow field data sets, and show the performance and scalability on HPC systems.

In-Situ Feature Extraction of Large Scale Combustion Simulations Using Segmented Merge Trees

Aaditya G. Landge, Valerio Pascucci, Attila Gyulassy (University of Utah), Janine C. Bennett, Hemanth Kolla, Jacqueline Chen (Sandia National Laboratories), Peer-Timo Bremer (Lawrence Livermore National Laboratory)

The increasing amount of data generated by scientific simulations coupled with system I/O constraints is fueling a need for in-situ analysis techniques. Of particular interest are approaches that produce reduced data representations while maintaining the ability to redefine, extract, and study features in a post-process to obtain scientific insights. We present two variants of in-situ feature extraction techniques using segmented merge trees, which encode a wide range of threshold based features. The first approach is a fast, low communication cost technique that generates an exact solution but has limited scalability. The second is a scalable, local approximation that guarantees to correctly extract all features up to a predefined size. We demonstrate both variants using some of the largest combustion simulations available on leadership class supercomputers. Our approach allows state-of-the-art, feature-based analysis to be performed in-situ at significantly higher frequency than currently possible and with negligible impact on the simulation runtime.

Memory System Energy Efficiency

Chair: Alex Ramirez (Polytechnic University of Catalonia)

3:30pm-5pm

Room: 393-94-95

ECC Parity: A Technique for Efficient Memory Error Resilience for Multi-Channel Memory Systems

Xun Jian, Rakesh Kumar (University of Illinois at Urbana-Champaign)

Servers and HPC systems often use a strong memory error correction code, or ECC, to meet their reliability and availability requirements. However, these ECCs often require significant capacity and/or power overheads. We observe that since memory channels are independent from one another, error correction only needs to be performed for one channel at a time. Based on this observation, we show that instead of always storing in memory the actual ECC correction bits as do existing systems, it is sufficient to store the bitwise parity of the ECC correction bits of different channels for fault-free memory regions, and store the actual ECC correction bits only for faulty regions. By trading off the resultant ECC capacity overhead reduction for improved memory energy efficiency, the proposed technique reduces memory energy per instruction by 54.4% and 18.5%, respectively, compared to commercial chipkill correct and DIMM-kill correct, while incurring similar or lower capacity overheads.

Using an Adaptive HPC Runtime System to Reconfigure the Cache Hierarchy

Ehsan Totoni, Josep Torrellas, Laxmikant V. Kale (University of Illinois at Urbana-Champaign)

A large portion (40% or more) of a processor's power and energy is consumed by the cache hierarchy. We propose a software-controlled adaptive runtime system-based reconfiguration approach for common HPC applications to save cache energy. Our approach overcomes the two major limitations associated with other methods that turn off ways of set-associative caches: predicting the application's future, and finding the best cache hierarchy configuration. Our approach uses Formal Language Theory to recognize the application's pattern and predict its future. Furthermore, it uses the prevalent Single Program Multiple Data (SPMD) model of HPC codes to find the best configuration in parallel quickly. Our experiments using cycle-accurate simulations indicate that 67% of cache energy can be saved by paying just 2.4% performance penalty on average. Moreover, we demonstrate that for some applications, switching to a software-controlled reconfigurable streaming strategy can improve performance by up to 30% and save 75% of cache energy.

Microbank: Architecting Through-Silicon Interposer-Based Main Memory Systems

Young Hoon Son, Seongil O (Seoul National University), Hyunggyun Yang (Pohang University of Science and Technology), Daejin Jung, Jung Ho Ahn (Seoul National University), John Kim (Korea Advanced Institute of Science and Technology), Jangwoo Kim (Pohang University of Science and Technology), Jae W. Lee (Sungkyunkwan University)

Through-Silicon Interposer (TSI) has been recently proposed to provide high memory bandwidth and improve energy efficiency of the main memory system. However, the impact of TSI on the main memory system architecture has not been well explored. While TSI improves the I/O energy efficiency, we show that it results in an unbalanced memory system design in terms of energy efficiency as the core DRAM dominates overall energy consumption. To balance and improve the energy efficiency of a TSI-based memory system, we propose microbank, a novel DRAM device organization in which each bank is partitioned into multiple smaller banks (called microbanks) that operate independently like conventional banks with minimal area overhead. The microbank organization significantly increases the amount of bank-level parallelism to improve performance while increasing energy efficiency of TSI-based memory system. The massive number of microbanks also simplifies the memory system design with the larger number of open DRAM rows.

Posters

Posters at SC14 provide an excellent opportunity for short presentations and informal discussions with conference attendees. SC14 posters display cutting-edge, interesting research in high performance computing, storage, networking and analytics.

In addition to regular posters, SC14 hosts the ACM Student Research Competition (SRC), where 23 student-developed posters selected after a rigorous peer-review process from 44 submissions are on display and compete for the best graduate and undergraduate student poster awards. The student posters are included in the Poster Exhibit display area and the poster presentations take place during the Poster Reception on Tuesday from 5:15pm to 7pm. During the Poster Reception, a jury selects posters for the ACM SRC semi-final. Selected students present their work on Wednesday, Nov. 19, at 3:30 pm in Room 404 of the convention center. Each selected student presents a 10-minute talk about her/his poster and key contributions. Please also join these finalists and encourage our next generation of SC researchers. A jury evaluates the students and chooses the winners for the ACM SRC graduate and undergraduate medals and cash prizes. Winners are presented at the SC14 award ceremony on Thursday. First place winners of the ACM SRC at SC14 are entered into the SRC Grand Finals (Grand Finals are judged over the Internet). Winners of the Grand Finals (with their faculty advisor) are invited to the annual ACM Awards Banquet where the ACM Turing Award is presented.

Posters

Tuesday, November 18

Poster Display (including ACM Student Research Competition Posters)

8:30am-7pm

Reception & Display (including ACM Student Research Competition Posters)

5:15pm-7pm

Room: New Orleans Theater Lobby

The Posters display, consisting of Research Posters and ACM Student Research Competition Posters posters kicks off at 8:30am, Tuesday, November 19, with a reception at 5:15pm. The reception is an opportunity for attendees to interact with poster presenters. The reception is open to all attendees with Technical Program registration. Complimentary refreshments and appetizers are available until 7pm.

Wednesday, November 19

Poster Display (including ACM Student Research Competition Posters)

8:30am-5pm

Room: New Orleans Theater Lobby

ACM Student Research Competition Presentations

Chair: Rob Schreiber (HP Labs)

3:30pm-5pm

Room: 386-87

Twenty-three student-developed posters are on display and compete for the best graduate and undergraduate student poster awards. The student posters are included in the Poster Exhibit display area. Finalists in the student poster competition present their work during this time, and the winners are announced at the award ceremony. The first place graduate and undergraduate students will go on to compete in the 2015 ACM SRC Grand Finals. Posters will remain on display through Thursday, November 20.

Thursday, November 20

Poster Display (including ACM Student Research Competition Posters)

8:30am-5pm

Room: New Orleans Theater Lobby

DySectAPI: A Massively Scalable Prescription-Based Debugging Model

Nicklas Bo Jensen, Niklas Quarfot Nielsen, Sven Karlsson (Technical University of Denmark), Gregory L. Lee, Dong H. Ahn, Matthew Legendre, Martin Schulz (Lawrence Livermore National Laboratory)

We present DySectAPI, a tool that enables users to construct probe trees for automatic, event-driven debugging at scale. The traditional, interactive debugging model, whereby users manually step through and inspect their application, does not scale well even for current supercomputers. While lightweight debugging models scale well, they can currently only debug a subset of bug classes. DySectAPI fills the gap between these two approaches with a novel user-guided approach. Using both experimental results and analytical modeling we show that DySectAPI scales well and can run with a low overhead on current systems.

A Power API for the HPC Community

David DeBonis, Ryan E. Grant, Stephen L. Olivie, Michael Levenhagen, Suzanne M. Kelly, Kevin T. Pedretti, James H. Laros (Sandia National Laboratories)

Power measurement and control are necessary for the proper operation of future power capped HPC systems. Current approaches to power measurement and control are limited and have proprietary interfaces. This poster presents a new vendor-neutral API for power measurement and control, based on the API document developed at Sandia National Labs and reviewed by industry, laboratory, and university partners. The API provides a publicly available free interface design that can be used to leverage power measurement and control hardware for large systems. This poster presents a description of the API itself, and the motivation behind the design of the API and its individual interfaces. Finally, the poster summarizes several power related research projects at Sandia National Laboratories that can leverage the Power API to transition research efforts quickly and portably into a production domain.

Creating High-Performance Linear Algebra Using Lighthouse and Build-to-Order BLAS

Jeffrey J. Cook, Elizabeth R. Jessup (University of Colorado),
Sa-Lin C. Bernstein (Argonne National Laboratory), Boyana
Norris (University of Oregon)

Creating efficient linear algebra software is essential to HPC, but it is a time-consuming task. Typically, scientists use Basic Linear Algebra Subprograms (BLAS) as the mathematical building-blocks for complete algorithms. They can present a steep learning curve, however, and require extensive knowledge to optimize. In order to make BLAS and derivative libraries such as LAPACK more accessible, we present Lighthouse, a web-based taxonomy of dense and sparse linear system and eigenvalue solvers. With its new interface to the Build-to-Order BLAS compiler (BTO), which generates custom linear algebra kernels from MATLAB-like equations, it is easy for users to create high-performance code.

Scaling OpenMP Programs to a Thousand Cores on the Numascale Architecture

Dirk Schmidl (RWTH Aachen University), Atle Vesterkjær
(Numascale)

The downside of shared memory programming compared to message passing is the limitation to run on a single system. Numascale's interconnect couples several standard servers in a cache coherent way which allows shared memory programming on the complete machine. However, this does not necessarily mean that shared memory programs deliver satisfying performance on such a system. Here we investigated a Numascale machine with 1728 cores hosted at the University of Oslo. The system consists of 72 IBM x3755 M3 nodes coupled in a 3D torus network topology with Numascale's interconnect. We investigate the memory bandwidth with kernel benchmarks and furthermore look at an application developed at the Institute of Combustion Technology at RWTH Aachen University. We describe tuning steps to optimize the application for large SMP machines like the Numascale machine and present good performance results for OpenMP runs with 1024 threads on the Oslo system.

Monetary Cost Optimizations for HPC Applications on Amazon Clouds: Checkpoints and Replicated Execution

Yifan Gong, Bingsheng He (Nanyang Technological University)

Due to the emerging of MPI-based HPC or scientific applications in the cloud and the pay-as-you-go pricing of the cloud, monetary cost becomes one important optimization factor. In this study, we focus on monetary cost optimizations for running MPI-based applications with deadline constraints on Amazon EC2 clouds. We propose a runtime system called SOMPI

that minimizes the monetary cost and has the same interfaces as MPI. We leverage the hybrid of both spot and on-demand instances to reduce monetary cost. As a spot instance can fail at any time due to out-of-bid events, fault tolerant executions are necessary. We propose a cost model to guide the combination of two common fault-tolerant mechanisms, including checkpoint and replication, to have more cost-effective reliable executions. The experimental results with NPB benchmarks on Amazon EC2 demonstrate (1) the significant monetary cost reduction and performance improvement; (2) the necessity of combining checkpoint and replication techniques.

Rolls-Royce Hydra CFD Code on GPUs Using OP2 Abstraction

Istvan Z. Reguly, Gihan R. Mudalige (University of Oxford),
Carlo Bertolli (IBM), Michael B. Giles (University of Oxford),
Adam Betts, Paul H. J. Kelly (Imperial College London),
David Radford (Rolls-Royce plc)

Hydra is an industrial CFD application used for the design of turbomachinery, now automatically accelerated by GPUs through the OP2 domain specific "active library" for unstructured grid algorithms. From the high-level definition, either CPU or GPU code is generated, applying optimizations such as conversion to Structure-of-Arrays, use of the read-only cache or the tuning of block sizes automatically. A single GPU is over two times faster than the original on a server-class CPU. We demonstrate excellent strong and weak scaling, using up to 4096 CPU cores or 16 GPUs.

10x Acceleration of Fusion Plasma Edge Simulations Using the Parareal Algorithm

Debasmita Samaddar (Culham Centre for Fusion Energy/UK Atomic Energy Authority), David P. Coster (Max Planck Institute of Plasma Physics), Xavier Bonnin (National Center for Scientific Research - Science Laboratory Methods and Materials), Eva Havlickova (Culham Centre for Fusion Energy/UK Atomic Energy Authority), Wael R. Elwasif, Lee A. Berry, Donald B. Batchelor (Oak Ridge National Laboratory)

Simulations involving the edge of magnetically confined, fusion plasma are an extremely challenging task. The physics in this regime is particularly complex due to the presence of neutrals and the device wall. These simulations are extremely computationally intensive but are key to rapidly achieving thermonuclear breakeven on ITER-like machines. Space parallelization has hitherto been the most common approach to the best of our knowledge. Parallelizing time adds a new dimension to the parallelization thus significantly increasing computational speedup and resource utilization. This poster describes the application of the parareal algorithm to the edge-plasma code package – SOLPS. The algorithm requires a coarse (G) and a fine (F) solver – and optimizing G is a challenge. This work

explores two approaches for G leading to computational gain >10 which is significant for simulations of this kind. It is also shown that an event-based approach to the algorithm greatly enhances performance.

Raexplore: Enabling Rapid, Automated Architecture Exploration for Full Applications

Yao Zhang, Prasanna Balaprakash, Jiayuan Meng, Vitali Morozov, Scott Parker, Kalyan Kumaran (Argonne National Laboratory)

We present Raexplore, a performance modeling framework for architecture exploration. Raexplore enables rapid, automated, and systematic search of architecture design space by combining hardware counter-based performance characterization and analytical performance modeling. We demonstrate Raexplore for two recent manycore processors IBM Blue-Gene/Q compute chip and Intel Xeon Phi, targeting a set of scientific applications. Our framework is able to capture complex interactions between architectural components including instruction pipeline, cache, and memory, and to achieve a 3–22% error for same-architecture and cross-architecture performance predictions. Furthermore, we apply our framework to assess the two processors, and discover and evaluate a list of potential architectural scaling options for future processor designs.

Parallel Clustering Coefficient Computation Using GPUs

Tahsin Reza (University of British Columbia), Tanuj Kr Aasawat (Jadavpur University), Matei Ripeanu (University of British Columbia)

Clustering coefficient is the measure of how tightly vertices are bounded in a network. The Triangle Counting problem is at the core of clustering coefficient computation. We present a new technique for implementing clustering coefficient algorithm on GPUs. It relies on neighbor lists being sorted with respect to vertex ID. The algorithm can process very large graphs not previously seen in the literature for single-node in-memory systems. Our technique is able to compute a clustering coefficient of each vertex in power-law graphs with up to 512M edges on a single GPU. Our GPU implementation offers seven times speed up over the best known work for the same graph. For the CPU implementation, we present results for graphs with up to 4B edges. We also investigate performance bottleneck of power-law graphs with skewed vertex degree distribution by analyzing performance counter outputs and interpreting them in terms of GPU memory and thread characteristics.

PACC: An Extension of OpenACC for Pipelined Processing of Large Data on a GPU

Tomochika Kato, Fumihiko Ino, Kenichi Hagihara (Osaka University)

We present a suite of directives, named pipelined accelerator (PACC), and its implementation for accelerating large-scale computation on a graphics processing unit (GPU). PACC extends OpenACC to achieve division of large data that cannot be entirely stored in device memory. Given a program with PACC directives, our PACC translator rewrites the program into an OpenACC program such that data is divided into multiple chunks for accelerated execution. Furthermore, the generated program processes chunks in a pipeline so that data transfer between the CPU and GPU can overlap with computation on the GPU. Some preliminary results are also presented to show the impact of PACC in terms of the program execution time and the maximum data size that can be processed successfully.

Award: Best Poster Finalist

Scalable Arbitrary-Order Pseudo-Spectral Electromagnetic Solver

Jean-Luc Vay, Leroy Anthony Drummond, Alice Koniges (Lawrence Berkeley National Laboratory), Brendan Godfrey (University of Maryland and Lawrence Berkeley National Laboratory), Irving Haber (University of Maryland)

Numerical simulations have been critical in the recent rapid developments of advanced accelerator concepts. Among the various available numerical techniques, the Particle-In-Cell (PIC) approach is the method of choice for self-consistent simulations from first principles. While spectral methods have been popular in the early PIC codes, the finite-difference time-domain method has become dominant. Recently, a novel parallelization strategy was proposed that takes advantage of the local nature of Maxwell equations that has the potential to combine spectral accuracy with finite-difference favorable parallel scaling. Due to its compute-intensive nature combined with adjustable accuracy and locality, the new solver promises to be especially well suited for emerging exascale systems. The new solver was recently extended to enable user-programmability of the spatial and temporal order of accuracy at runtime, enabling a level of scalability and flexibility that is unprecedented for such codes.

Characterizing Application Sensitivity to Network Performance

Eli Rosenthal (Brown University), Edgar A. Leon (Lawrence Livermore National Laboratory)

As systems grow in scale, network performance becomes an important component of the time-to-solution of scientific, message-passing applications. We expect that future systems may include 100K nodes connected through a high-speed

interconnect fabric. In this work, we investigate the impact of multi-rail networking on the performance of a representative set of HPC mini-applications to answer the following questions: Are applications limited by network performance? Can they leverage improvements in network bandwidth? What type of network operations and message sizes do they use? This work provides a better understanding of the network requirements of applications and helps us determine the cost-value added by higher-speed networking solutions in the procurement of future systems.

HPC for Students Enrolled in Science and Non-Science Degree Programs

Suzanne McIntosh (New York University Courant Institute and Cloudera)

The unprecedented growth in applications that leverage high performance computing (HPC) platforms ranging from supercomputers and grid computing systems, to Hadoop clusters, has created demand for graduates with skills in parallel programming, scientific computing, data science, and big data engineering.

Traditionally, HPC courses have been part of the Computer Science curriculum. As HPC applications find their way into finance, medicine, life sciences, advertising, and other industries, an HPC curriculum to simultaneously address the needs of Computer Science students and students studying business, life sciences, or mathematics is required.

GPU Acceleration of Small Dense Matrix Computation of the One-Sided Factorizations

Tingxing Dong, Mark Gates, Azzam Haidar, Piotr Luszczek, Stanimire Tomov (University of Tennessee, Knoxville)

In scientific applications, one often needs to solve many small size problems. The size of each of these small linear systems depends, for example, on the number of the ordinary differential equations (ODEs) used in the model, and can be on the order of hundreds of unknowns. To efficiently exploit the computing power of modern accelerator hardware, these linear systems are processed in batches. The state-of-the-art libraries for linear algebra that target GPUs, such as MAGMA, focus on large matrix sizes. They change the data layout by transposing the matrix to avoid these divergence and non-coalescing penalties. However, the data movement associated with transposition is very expensive for small matrices. We propose a batched one-sided factorizations for GPUs by using a multi-level blocked right looking algorithm that preserves the data layout but minimizes the penalty of partial pivoting. Our implementation achieves many-fold speedup when compared to the alternatives.

Large-Scale Granular Simulations Using Dynamic Load Balance on a GPU Supercomputer

Satori Tsuzuki, Takayuki Aoki (Tokyo Institute of Technology)

A billion particles are required to describe granular phenomena by using particle simulations based on Distinct Element Method (DEM), which computes contact interactions among particles. Multiple GPUs on TSUBAME 2.5 in Tokyo Tech are used to boost the simulation and are assigned to each domain decomposed in space. Since the particle distribution changes in time and space, dynamic domain decomposition is required for large-scale DEM particle simulations. We have introduced a two-dimensional slice-grid method to keep the same number of particles for each domain. Due to particles crossing the sub-domain boundary, the memory used for active particles is fragmented and the optimum frequency of de-fragmentation on a CPU is studied by taking account of data transfer cost through the PCI-Express bus. Our dynamic load balancing works well with good scalability in proportion to the number of GPUs. We demonstrate several DEM simulations running on TSUBAME 2.5.

HPC Lessons from Fan-In Communications

Terry Jones, Bradley Settlemyer (Oak Ridge National Laboratory)

We present a study of an important class of communication operations—the fan-in communication pattern. By its nature, fan-in communications form ‘hot spots’ that present significant challenges for any interconnect fabric and communication software stack. Yet despite the inherent challenges, these communication patterns are common in both applications (which often perform reductions and other collective operations that include fan-in communication such as barriers) and system software (where they assume an important role within parallel file systems and other components requiring high-bandwidth or low-latency I/O). Our study determines the effectiveness of differing client-server fan-in strategies. We describe fan-in performance in terms of aggregate bandwidth in the presence of varying degrees of congestion, as well as several other key attributes. Results are presented for two representative supercomputers (a large Cray Aries-interconnect based super-computer and a large Gemini-interconnect based super-computer). Finally, we provide recommended communication strategies based on our findings.

Bandwidth-Aware Resource Management for Extreme Scale Systems

Zhou Zhou, Xu Yang, Zhiling Lan (Illinois Institute of Technology), Narayan Desai (Ericsson Inc.), Paul Rich, Vitali Morozov, Wei Tang (Argonne National Laboratory)

As systems scale towards exascale, many resources including the traditional CPU cycles as well as non-traditional resources

(e.g., communication bandwidth) will become increasingly constrained. This change will pose critical challenges on resource management and job scheduling. As systems continue to evolve, we expect non-traditional resources like communication bandwidth to increasingly be explicitly allocated, where they have previously been managed in an implicit fashion. In this paper we investigate smart allocation of communication bandwidth on Blue Gene systems. The partition-based design in Blue Gene systems provides us a unique opportunity to explicitly allocate bandwidth to jobs, in a way that isn't possible on other systems. While this capability is currently rare, we expect it to become more common in the future. This paper makes two major contributions. The first is substantial benchmarking of leadership applications, focusing on assessing application sensitivity to communication bandwidth at large scale.

Comparing Algorithms for Detecting Abrupt Change Points in Data

Cody L. Buntain (University of Maryland), Christopher Natoli (University of Chicago), Miroslav Živković (University of Amsterdam)

Detecting points in data where the underlying distribution changes is not a new task, but much of the existing literature assumes univariate and independent data, assumptions often violated in real data sets. This work addresses this gap in the literature by implementing a set of change point detection algorithms and a test harness for evaluating their performance and relative strengths and weaknesses in multi-variate data of varying dimension and temporal dependence. We then apply our implementations to real-world data taken from structural sensors placed on laboratory a bridge and two years of Bitcoin market data from the Mt. Gox exchange. Although more work is necessary to explore these real-world data sets more thoroughly, our results demonstrate circumstances in which an online, non-parametric algorithm does and does not perform as well as offline, parametric algorithms and provides an early foundation for future investigations.

Performance of Block Jacobi-Davidson Algorithms

Melven Roehrig-Zoellner, Jonas Thies (German Aerospace Center), Moritz Kreutzer (Erlangen Regional Computing Center), Andreas Alvermann, Andreas Pieper (University of Greifswald, Institute of Physics), Achim Basermann (German Aerospace Center), Georg Hager, Gerhard Wellein (Erlangen Regional Computing Center), Holger Fehske (University of Greifswald, Institute of Physics)

Jacobi-Davidson methods can efficiently compute a few eigenpairs of a large sparse matrix. Block variants of Jacobi-Davidson are known to be more robust than the standard algorithm, but they are usually avoided as the total number of floating point operations increases. We present the implementation of a block Jacobi-Davidson solver and show by detailed

performance engineering and numerical experiments that the increase in operations is typically more than compensated by performance gains on modern architectures, leading to a method that is both more efficient and robust than its single vector counterpart.

XcalableACC – a Directive-Based Language Extension for Accelerated Parallel Computing

Hitoshi Murai, Masahiro Nakao, Takenori Shimosaka (RIKEN), Akihiro Tabuchi, Taisuke Boku, Mitsuhisa Sato (University of Tsukuba)

Accelerated parallel computers (APC) such as GPU clusters are emerging as an HPC platform. This poster proposes Xcalable-ACC, a directive-based language extension to program APCs. XcalableACC is a combination of two existing directive-based languages, XcalableMP and OpenACC, and has two additional functions: data/work mapping among multiple accelerators and direct communication between accelerators. The result of the preliminary evaluation shows that XcalableACC is a promising means to program APCs. The basic concept of XcalableACC and the result of preliminary evaluation is presented.

Exploring Hybrid Hardware and Data Placement Strategies for the Graph 500 Challenge

Scott Sallinen, Daniel Borges, Abdullah Gharaibeh, Matei Ripeanu (University of British Columbia)

Our research presents our experience with exploring the configuration space and data placement strategies for Breadth First Search (BFS) on large-scale graphs in the context of a hybrid, GPU+CPU architecture and the Graph 500 challenge. Recent work on GPU graph traversal has often targeted relatively small graphs that fit on device memory only; our goal is to process large graphs that stretch the limits of single-node commodity machines. When processing large graphs with tens of billions of edges or more, the techniques we explore are crucial to obtain an optimal performance level on hybrid CPU+GPU architectures. In particular, we show that the following configuration options can significantly impact performance: (i) the partitioning strategy, (ii) the placement of the graph partitions in memory in a 'sorted by degree' order, and (iii) the strategy to manage memory used for the GPU partitions when the graph no longer fits in the device memory.

Using IKAROS to Form Scalable Storage Platforms

Christos Filippidis (National Center of Scientific Research Demokritos), Yiannis Cotronis (University of Athens), Christos Markou (National Center of Scientific Research Demokritos)

We present IKAROS as a utility that permits us to form virtual scalable storage formations in order to create more user-driven computing facilities with application users and owners playing a decisive role in governance and focusing on placing computer science and the harvesting of 'big data' at the center

of scientific discovery. IKAROS enables users to virtually connect the several different computing facilities (Grids, Clouds, HPCs, Data Centers, Local computing Clusters and personal storage devices), being utilized by an international collaborative experiment, on-demand based to their needs.

Performance of Sparse Matrix-Multiple Vectors Multiplication on Multicore and GPUs

Walid Abu-Sufah (University of Illinois at Urbana-Champaign and University of Jordan), Khalid Ahmad (University of Jordan)

Sparse matrix-vector and multiple-vector multiplications (SpMV and SpMM) are performance bottlenecks in numerous applications. We implemented two SpMM kernels to integrate in our library of auto-tuned kernels for GPUs. Our kernels use registers to exploit data reuse in SpMM. DIA-SpMM targets structured matrices and ELL-SpMM targets matrices with uniform row lengths. Work is continuing on SpMM kernels for unstructured matrices.

Executing on NVIDIA Kepler Tesla K40m, DIA-SpMM is 2.4x faster than NVIDIA CUSP DIA-SpMV. ELL-SpMM is 2.8x faster than CUSP ELL-SpMV; DIA-SpMM is 5.2x faster than the highly optimized NVIDIA CUSPARSE CSR-SpMV. The maximum speedup is 6.5x. ELL-SpMM is 3.9x faster than CUSPARSE CSR-SpMV. The maximum speedup is 8.3x; DIA-SpMM is 2x faster than CUSPARSE CSR-SpMM. ELL-SpMM is 1.6x faster; For structured matrices, DIA-SpMM on the K40m GPU is 7.2x faster than Intel MKL CSR-SpMV on a dual socket 10-core Intel Ivy Bridge E5-2690. The maximum speedup is 12.3x.

The MDTM Project

Liang Zhang (Fermi National Laboratory), Tan Li, Yufei Ren (Stony Brook University), Phil DeMar (Fermi National Laboratory), Shudong Jin (Stony Brook University), Dantong Yu (Brookhaven National Laboratory), Wenji Wu (Fermi National Laboratory)

Multicore and manycore have become the norm for high performance computing. These new architectures provide advanced features that can be exploited to design and implement a new generation of high-performance data movement tools. DOE ASCR program is funding FNAL and BNL to collaboratively work on a Multicore-Aware Data Transfer Middleware (MDTM) project. MDTM aims to accelerate data movement toolkits on multicore systems. Essentially, MDTM consists of two major components: MDTM middleware services to harness multicore parallelism and make intelligent decision to align CPU, memory, and I/O device together to expedite higher layer applications, and 2) MDTM data transfer applications (client or server) that utilize the middleware to reserve and manage multiple CPU cores, memory, network devices, disk storage as an integrated resource entity and achieve high

throughput and improved quality of servers compared with those existing approaches.

Hydra: An HTML5-Based Application for High-Throughput Visualization of Ligand Docking

Yuan Zhao (University of California, San Diego), Jason Haga (National Institute of Advanced Industrial Science & Technology)

Hydra is an HTML5-based application for high-throughput visualization of molecular docking simulations. Unlike existing solutions, Hydra's implementation is platform agnostic, and therefore can be deployed quickly and cheaply across various hardware configurations. Additionally, it is designed with an intuitive interface that is scalable with respect to screen sizes, ranging from mobile devices to large, tiled display walls (TDWs).

Scalable and Highly Available Fault Resilient Programming Middleware for Exascale Computing

Atsuko Takefusa, Tsutomu Ikegami, Hidemoto Nakada, Ryousei Takano, Takayuki Tozawa, Yoshio Tanaka (National Institute of Advanced Industrial Science & Technology)

Falanx is a programming middleware for the development of applications for exascale computing. Because of the fragility of the computing environment, applications are required to be not only scalable, but also fault resilient. Falanx employs an MPI-based hierarchical parallel programming model for the scalability. Falanx consists of Resource Management System (RMS) and Data Store (DS): The RMS allocates processes of each task to computing nodes avoiding failed nodes. The DS redundantly preserves data required for each application, and prevents data loss. It is necessary that these components must be scalable and that they themselves have to be implemented in a fault resilient manner. We design a scalable and highly available middleware, which consists of RMS and DS, and implement them by using Apache ZooKeeper and Kyoto Cabinet. Then, we investigate the basic performance from the preliminary experiments and confirm the feasibility from experiments using an actual application, OpenFMO.

I/O Monitoring in a Hadoop Cluster

Carson L. Wiens (Los Alamos National Laboratory), Joshua M. C. Long (Los Alamos National Laboratory), Joel R. Ornstein (Los Alamos National Laboratory)

As data sizes for scientific computations grow larger, more of these types of computations are bottlenecked by disk input and output, rather than processing speed. Monitoring reads and writes, therefore, has become an important component of distributed computing clusters. Our team investigated several different possibilities for monitoring I/O on a Hadoop cluster,

including the Splunk app for HadoopOps, Ganglia, and log file output. Each of these three methods was evaluated for compatibility, ease of use, and display. Surprisingly, despite the fact that Splunk HadoopOps is made specifically for Hadoop clusters, other monitoring programs and techniques still proved to be useful. Using our monitoring tools, we were also able to observe input and output behavior over different cluster architectures.

Fault Injection, Detection, and Correction in CLAMR Using F-SEFI

Brian Atkinson (Clemson University), Nathan DeBardeleben, Qiang Guan (Los Alamos National Laboratory)

F-SEFI is a fine-grained software-based soft fault injection tool developed at LANL. We used F-SEFI to study the resilience of the scientific application to CLAMR, a cell based adaptive mesh refinement hydrodynamic code also developed at LANL, in the presence of soft errors. CLAMR models a cylindrical shock generated in the center of the mesh that reflects off the boundaries. We focused our fault injections on the floating point add operations in the exponent bit field. Using conservation of mass calculations inherent to the shallow water simulations, we specified an acceptable bound for the mass percentage difference between specified time steps. We built a checkpointing and rollback mechanisms into CLAMR to save and restore state and mesh values from backup files. Using the checkpointing and roll back routines, we were able to recover from 81% of soft errors that would have caused incorrect results or the application to crash.

Massively Parallel and Near Linear Time Graph Analytics

Fazle Elahi Faisal, Yves Ineichen, A. Cristiano I. Malossi, Peter Staar, Costas Bekas (IBM Zurich Research Laboratory), Alessandro Curioni (IBM Zurich Research Laboratory)

Graph theory and graph analytics are essential in many research problems and have applications in various domains. The rapid growth in the data can lead to very large sparse graphs consisting of billions of nodes. This poses many challenges in designing fast algorithms for large-scale graph analytics.

We believe that in the era of the big data regime accuracy has to be traded for algorithmic complexity and scalability to drive big data analytics. We discuss the underlying mathematical models of two novel near linear ($O(N)$) methods for graph analytics: estimating subgraph centralities, and spectrograms. Parallelism on multiple levels is used to attain good scalability. We benchmarked the proposed algorithms on huge datasets, e.g. the European street network containing 51 million nodes and 108 million non-zeros, on up to 32,768 cores. Due to the

iterative behavior of the methods, a very good estimate can be computed in a matter of seconds.

Big Data Analytics on Object Stores: A Performance Study

Lukas Rupperecht (Imperial College London), Rui Zhang, Dean Hildebrand (IBM Research - Almaden)

Object Stores provide a cheap solution for storing, sharing, and accessing globally distributed, scientific datasets. Their scalability and high reliability makes them well suited as archive storage for data produced by HPC clusters. To make this archive active, i.e. to process the archived data, a separate analytics cluster is needed which adds overhead in terms of cost, maintenance, and performance as data needs to be copied between two systems. Running the analytics directly on the Object Store can greatly simplify the archive usage. We study the problems that arise when running an analytics system (Hadoop MapReduce) on top of an Object Store (OpenStack Swift). We conduct a set of detailed micro- and macro benchmarks and find that the static, consistent hashing based, object-node mapping adds significant overhead when object writes are not local. Additionally, repeated authentication calls slow down jobs which have a high performance impact, especially on interactive workloads.

Using Global View Resilience (GVR) to add Resilience to Exascale Applications

Hajime Fujita, Nan Dun (University of Chicago/Argonne National Laboratory), Aiman Fang (University of Chicago), Zachary A. Rubenstein (University of Chicago), Ziming Zheng (HP Vertica), Kamil Iskra (Argonne National Laboratory), Jeff Hammond (Intel Corporation), Anshu Dubey (Lawrence Berkeley National Laboratory), Pavan Balaji (Argonne National Laboratory), Andrew A. Chien (University of Chicago/Argonne National Laboratory)

Resilience is a big challenge in future exascale machines. Existing approaches are unlikely to address complex failures like latent errors; therefore, we need a new approach. We propose Global View Resilience (GVR), a new library that exploits a global view data model and adds reliability through versioning (multi-version), user control timing and rate (multi-stream), and flexible cross layer error signaling and recovery. GVR enables application programmers to exploit deep scientific and application code insights to manage resilience (and its overhead) in a flexible, portable fashion. We applied the GVR library to several existing scientific application codes and showed that GVR can be easily applied and runtime overhead for versioning is negligible.

Award: Best Poster Finalist

A Cloud-Based Interactive Data Infrastructure for Sensor Networks

Tonglin Li (Illinois Institute of Technology), Kate Keahey (Argonne National Laboratory), Ioan Raicu (Illinois Institute of Technology)

As small, sensor devices become ubiquitous, reliable, and cheaper, more domain sciences are creating “instruments at large” – dynamic, self-organizing, groups of sensors. This calls for an infrastructure able to collect, store, query, and process data set from sensor networks. There are some challenges. First is the need to interact with and administer the sensors remotely. The network connectivity for sensors is intermittent. This calls for communication that can withstand unreliable networks. In this work we present a set of protocols and a cloud-based data store called “WaggleDB” that address those challenges. The system aggregates and stores data from sensor networks and enables the users to query those data sets. It address the challenges above with a scalable multi-tier architecture, which is designed in such way that each tier can be scaled by adding more independent resources provisioned on-demand in the cloud.

Exposing MPI Objects for Debugging

Laust Brock-Nannestad (Technical University of Denmark), John DeSignore (Rogue Wave Software, Inc.), Jeffrey M. Squyres (Cisco Systems), Sven Karlsson (Technical University of Denmark), Kathryn Mohror (Lawrence Livermore National Laboratory)

Debugging is the art of inspecting a program’s state in search of incorrect behavior. In message passing applications a central part of this state lies within the MPI library. There’s no standardized interface to this information; at best a developer can manually inspect the underlying data structures and extract information. For closed-source MPI implementations that’s not even an option.

The MPI Tools Working Group has proposed a standard for MPI Handle Introspection. This defines an interface that—independently of MPI implementation—allows a debugger to extract information concerning MPI objects from an application, and present it to the developer in an understandable format.

We implement support for introspection in the TotalView debugger and test it against a reference introspection implementation in Open MPI. We explain how the debugger interfaces with the MPI implementation and demonstrate how introspection raises the abstraction level to simplify debugging of MPI programming errors.

MetaMorph: A Modular Library for Democratizing the Acceleration of Parallel Computing Across Heterogeneous Devices

Paul D. Sathre (Virginia Polytechnic Institute and State University), Wu-chun Feng (Virginia Polytechnic Institute and State University)

Heterogeneity in computing has become ubiquitous. Computing systems ranging from smartphones to supercomputers now consist of multiple types of computing brains, typically at least one multi-core CPU and a many-core GPU. Such systems offer the promise of increased performance and energy efficiency over CPUs if the resources within such systems are used judiciously.

Alas, extracting the full performance potential from accelerator devices requires architectural expertise, expertise that is in short supply. Thus, there exists a need for software tools and ecosystems to support the development, maintenance, and upgrading of accelerated codes by non-expert domain scientists. To address this need, we present our initial prototype of a modular library of drop-in accelerated functions, which abstracts current and future accelerator backends behind a single, unified API, known as MetaMorph. By incrementally replacing computational primitives with calls to the MetaMorph API, domain scientists can transparently migrate code to current and future accelerator devices.

Skeptical Programming and Selective Reliability

James Elliott (North Carolina State University), Mark Hoemmen (Sandia National Laboratories), Frank Mueller (North Carolina State University)

Problem: Current and future systems may experience transient faults that silently corrupt data being operated on. These faults are troubling, because there is no indication a fault occurred and the cost to detect an error may involve performing operations multiple times and voting to determine if any values are tainted.

Approach: I focus on numerical methods and have proposed an approach called Skeptical Programming. My research couples the numerical method and the properties of data being operated on to derive cheap invariants that filter out large “damaging” errors, while allowing small “bounded” errors to slip through. The errors that slip through are easily handled by convergence theory, resulting in a low-overhead algorithm-based fault tolerance approach that exploits both numerical analysis and system-level fault tolerance. This technique scales well since we do not require additional communications and is applicable to many methods (we present findings for the CG and GMRES solvers).

Award: Best Poster Finalist

A High Performance C++ Generic Benchmark for Computational Epidemiology

Aniket Pugaonkar, Sandeep Gupta, Keith R. Bisset, Madhav V. Marathe (Virginia Polytechnic Institute and State University)

We design a benchmark consisting of several kernels which capture the essential compute, communication, and, data access patterns of high performance contagion-diffusion simulations used in computational networked epidemiology. The goal is to (a) derive different implementations for computing the contagion by combining different implementation of the kernels, and (b) evaluate which combination of implementation, run-time, and, hardware is most effective in running large scale contagion diffusion simulations.

Our proposed benchmark is designed using C++ generic programming primitives and lifting sequential strategies for parallel computations. Together these lead to a succinct description of the benchmark and significant code re-use when deriving strategies for new hardware. These aspects are crucial for an effective benchmark because the potential combination of hardware and runtimes are growing rapidly thereby making infeasible to write optimized strategy for the complete contagion diffusion from ground up for each compute system.

A Framework for Resource Aware Multithreading

Sunil Shrestha (University of Delaware), Joseph Manzano, Andres Marquez (Pacific Northwest National Laboratory), Guang Gao (University of Delaware)

In this poster, we present a novel methodology that takes into consideration multithreaded many-core designs to increase both intra and inter tile parallelism as well as memory residence on tilable applications. It partly takes advantage of polyhedral analysis and transformation, combined with a highly optimized fine grain tile runtime to exploit parallelism at multiple levels of the memory hierarchy. The main contributions include (1) a framework for multi-hierarchical tiling techniques that takes advantage of intra-tile parallel-start parallelism, (2) a data-flow inspired runtime library that allows the expression of parallelism with an efficient synchronization registry, and (3) an efficient memory reuse strategy. Our current implementation shows significant performance improvements on an Intel Xeon Phi board against instances produced by state-of-the-art compiler frameworks for selected loop nests.

Award: Best Poster Finalist

A Framework for Analyzing the Community Land Model within the Community Earth System Models

Dali Wang (Oak Ridge National Laboratory), Wei Wu, Yang Xu (University of Tennessee), Tomislav Janjusic, Wei Ding, Frank Winkler, Nick Forrington, Oscar Hernandez (Oak Ridge National Laboratory)

As environmental models (such as Accelerated Climate Model for Energy (ACME), Parallel Reactive Flow and Transport Model (PFLOTTRAN), Arctic Terrestrial Simulator (ATS), etc.) become more and more complicated, we need new tools to expedite integrated model development and facilitate the collaborations between field scientists, environmental system modelers and computer scientists. In this poster, we present our methods and efforts to analyze the Community Land Model (CLM), a terrestrial ecosystem model within the Community Earth System Models (CESM). Specifically, we demonstrate our objectives, methods and software tools to support interactive software structure exploration and automatic functional testing code generation, compiler analysis and interesting works on code porting preparations for pre-exascale computers. We believe that our experience on the environmental model, CLM, can be beneficial to many other scientific research programs which adapt the integrated, component-based modeling methodology on high-end computers.

Advanced Computation for High Intensity Accelerators

James Amundson, Qiming Lu (Fermi National Laboratory)

The ComPASS project is a US DOE SciDAC-funded collaboration of national labs, universities and industrial partners dedicated to creating particle accelerator simulation software. The future of particle physics requires high-intensity accelerators, whose planning, construction and operation require high fidelity simulations. We describe the methods and algorithms the ComPASS collaboration has employed to produce software capable of efficiently utilizing today's and tomorrow's high performance computing hardware. We also present the results of some key simulations and show how they are advancing accelerator science.

Parallel High-Order Geometric Multigrid Methods on Adaptive Meshes for Highly Heterogeneous Nonlinear Stokes Flow Simulations of Earth's Mantle

Johann Rudi (University of Texas at Austin), Hari Sundar (University of Utah), Tobin Isaac, Georg Stadler (University of Texas at Austin), Michael Gurnis (California Institute of Technology), Omar Ghattas (University of Texas at Austin)

The simulation of Earth's mantle flow with associated plate motion at global scale is challenging due to (1) the severe non-linear rheology, (2) the viscosity variations, and (3) the widely-varying spatial scales inherent in the problem. We discretize the governing nonlinear incompressible Stokes equations

using high-order finite elements on highly adapted meshes, which allow us to resolve plate boundaries down to a few hundred meters. Crucial components in our scalable solver are an efficient, collective communication-free, parallel implementation of geometric multigrid for high-order elements, the use of a Schur complement preconditioner that is robust with respect to extreme viscosity variations, and the use of an inexact Newton solver combined with grid-continuation methods. We present results based on real Earth data that are likely the most realistic global scale mantle flow simulations conducted to date.

Award: Best Poster Finalist

Visualizing the Behavior of Large Programs

Hoa Nguyen (University of Utah), Greg Bronevetsky (Lawrence Livermore National Laboratory)

As HPC simulations grow more complex, debugging and optimizing them becomes vastly more difficult. This makes it necessary to design visualization tools that help developers understand the behavior of their applications. We propose novel visualizations and effective mechanisms to interact with them to support developer efforts to easily and efficiently analyze the behavior of large applications. These visualization techniques address some of the major issues still plaguing program behavior analysis, including reducing human error, improving speed and quality of decision making, and mitigating program size and complexity issues in analyzing behaviors. A case study on HPC programs demonstrates that our methods can help developers efficiently explore the behavior of their applications.

Optimizing CAD and Mesh Generation Workflow for SeisSol

Sebastian Rettenberger (Technical University Munich), Cameron Smith (Rensselaer Polytechnic Institute), Christian Pelties (Ludwig Maximilians University Munich)

For accurate modeling of realistic earthquake scenarios it is very important to capture small-scale geometric features. In this poster, we optimized the whole pipeline from the different CAD sources (e.g. topography, fault structure) to a fully adaptive, unstructured tetrahedral mesh that can be used in SeisSol for simulating detailed rupture processes and seismic wave propagations at petascale performance. We present an automated CAD model generation which can easily take weeks or even month when done manually. From the CAD model a parallel meshing software - in the former workflow only serial mesh generators were able to write appropriate mesh files - generates an unstructured mesh with billions of tetrahedrons that preserves the details of the model. Due to our new mesh format, we are now able to initialize meshes at this size within seconds in SeisSol.

Diagnosing the Causes and Severity of One-Sided Message Contention

Nathan R. Tallent, Abhinav Vishnu, Hubertus Van Dam, Jeff Daily, Darren Kerbyson, Adolfo Hoisie (Pacific Northwest National Laboratory)

Two trends suggest network contention for one-sided messages is poised to become a performance problem that concerns application developers: an increased interest in one-sided programming models and a rising ratio of hardware threads to network injection bandwidth. We present effective and portable techniques for diagnosing the causes and severity of one-sided message contention. To detect that a message is affected by contention, we maintain statistics representing instantaneous (non-local) network resource demand. Using lightweight measurement and modeling, we identify the portion of a message's latency that is due to contention and whether contention occurs at the initiator or target. We attribute these metrics to program statements in their full static and dynamic context. Our results pinpoint the sources of contention, estimate their severity, and show that when message delivery time deviates from an ideal model, there are other messages contending for the same network links.

Award: Best Poster Finalist

Optimization and Highly Parallel Implementation of Domain Decomposition Based Algorithms

Lubomir Riha, Tomas Brzobohaty, Alexandros Markopoulos, Marta Jarosova, Tomas Kozubek (IT4Innovations)

We describe an implementation and scalability results of a hybrid FETI (Finite Element Tearing and Interconnecting) solver based on our variant of the FETI type domain decomposition method called Total FETI. In our approach a small number of neighboring subdomains is aggregated into clusters, which results into a smaller coarse problem. Current implementation of the solver is focused on the optimal performance of the main CG solver, including: implementation of communication hiding and avoiding techniques for global communications; optimization of the nearest neighbor communication - multiplication with global gluing matrix; and optimization of the parallel CG algorithm to iterate over local Lagrange multipliers only. The performance is demonstrated on a linear elasticity synthetic 3D cube and real world benchmarks.

Lightweight Scheduling for Balancing the Tradeoff Between Load Balance and Locality

Vivek Kale (University of Illinois), William Gropp (University of Illinois at Urbana-Champaign)

Many factors contribute to performance irregularities on massively parallel processors, leading to a load imbalance and loss of performance. At the same time, shared memory, multicore nodes make it easier to redistribute work within a node, particularly for work expressed in parallel do/for loops.

How can work be distributed across cores without disturbing data locality important to HPC codes? Additionally, how much better can we do than well-known strategies such as “guided” scheduling? We develop scheduling strategies and tuning mechanisms to find the right balance between load balance and locality. We also develop a basic runtime system and library to minimize programmer effort in applying these strategies. Our techniques provide 25.13% performance gains over static scheduling and 12.62% gains over guided scheduling for a widely used regular mesh benchmark at scale, and 38.16% performance gains over static scheduling and 20.45% gains over guided scheduling for a galaxy simulation at scale.

GPGPU-Enabled HPC Cloud Platform Based on OpenStack

Tae Joon Jun, Van Quoc Dung, Myong Hwan Yoo, Daeyoung Kim, HyeYoung Cho, Jaegyeon Hahm (Korea Institute of Science and Technology Information)

Currently, HPC users are interested in cloud computing as an alternative to supercomputers, and General Purpose Graphical Processing Units (GPGPUs) now play a crucial role in HPC. To provide GPGPU in cloud computing for HPC, we suggest a GPGPU HPC Cloud platform based on OpenStack. We combined OpenStack, KVM, and rCUDA to enable scalable use of GPUs among virtual machines. To use GPU resources more efficiently, we propose centralized, distributed GPU resource schedulers on cloud computing, mainly considering GPU resource usage, cloud network topology and traffic. Finally, we evaluated performance of our GPGPU HPC cloud platform. Compared to a cloud platform with PCI pass-through, speedup ratio was 0.93 to 0.99 on compute-intensive programs and 0.26 to 0.89 on data intensive programs.

Optimizing Stencil Computations: Multicore-Optimized Wavefront Diamond Blocking on Shared and Distributed Memory Systems

Tareq Malas (King Abdullah University of Science and Technology), Georg Hager (Erlangen Regional Computing Center), Hatem Ltaief (King Abdullah University of Science & Technology), Holger Stengel (University of Erlangen-Nuremberg), Gerhard Wellein (University of Erlangen-Nuremberg), David Keyes (King Abdullah University of Science & Technology)

The importance of stencil-based algorithms in computational science has focused attention on optimized parallel implementations for modern multilevel cache-based processors. Temporal blocking schemes leverage the large bandwidth and low latency of caches to accelerate stencil updates and approach theoretical peak performance. A key ingredient is the reduction of data traffic across slow data paths, especially the main memory interface. In this work we combine the ideas of multicore wavefront temporal blocking and diamond tiling to arrive

at generic stencil update schemes that show large reductions in memory pressure compared to conventional approaches. The resulting update schemes show performance advantages especially in bandwidth-starved situations, where low memory bandwidth or variable coefficients lead to strong saturation effects for non-temporally blocked code. We present performance results on modern Intel processors and apply advanced performance modeling techniques to reconcile the observed performance with hardware capabilities.

Efficient Data Compression by Efficient Use of HDF5 Format

Katsumi Hagita (National Defense Academy), Manabu Omiya (Hokkaido University), Takashi Honda (ZEON Corporation), Masao Ogino (Nagoya University)

Recently, in the area of supercomputing, data transfer problems have become significant. We proposed efficient data compression methods without changing common APIs like HDF5. JHPCN-DF (Jointed Hierarchical Precision Compression Number-Data Format) is a compression method using Bit Segmentation and Huffman coding. For visualizations and analyses, the lower bits of IEEE 754 format is not necessary. The required bits can be determined by direct check to numerical precisions corresponding to visualizations and analyses. We evaluated data size of HDF5 files coordinated by the JHPCN-DF framework in various simulations such as plasma electromagnetic particle-in-cell simulations, finite element method, and phase separated polymer materials. We confirmed visualization software works well with HDF5 files coordinated by the JHPCN-DF framework. For visualization, data sizes of HDF5 file seem to be reduced to 1/5 of original HDF5 file size. It is found that JHPCN-DF is effective for wide areas of simulations and data sciences.

Towards an Exascale implementation of Albany: a Performance-Portable Finite Element Application

Irina Demeshko (Sandia National Laboratories), H. Carter Edwards (Sandia National Laboratories), Michael A. Heroux (Sandia National Laboratories), Roger P. Pawlowski (Sandia National Laboratories), Eric T. Phipps (Sandia National Laboratories), Andrew G. Salinger (Sandia National Laboratories)

Modern HPC applications need to be run on many different platforms and performance portability has become a critical issue: parallel code needs to be executed correctly and performant despite variation in the architecture, operating system and software libraries. The numerical solution of partial differential equations using the finite element method is one of the key applications of high performance computing. This poster presents a performance portable implementation of the finite element assembly in the Albany code. This implementation is based on the Kokkos programming model from Trilinos, which

uses a library approach to provide performance portability across diverse devices with different memory models. Evaluation experiments show good performance results for a single implementation across three multicore/many-core architectures: NVIDIA GPUs, Multicore CPUs, Intel Xeon Phi.

Development of Distributed Parallel Explicit Moving Particle Simulation (MPS) Method and Zoom Up Tsunami Analysis on Urban Areas

Kohei Murotani, Seiichi Koshizuka (University of Tokyo), Masao Ogino (Nagoya University), Ryuji Shioya, Yasushi Nakabayashi (Tokyo University)

In this research, a distributed memory parallel algorithm of the explicit MPS (Moving Particle Simulation) method is described. The MPS method is one of the popular particle methods with collision. The ParMETIS is adopted for domain decomposition. We show the algorithm and the results of parallel scalability in our poster using the FX10 of the University of Tokyo. As applications, a large-scale run-up tsunami analysis such as inundating the Ishinomaki urban area and carrying two 10m diameter tanks by the tsunami is presented.

Early Evaluation of the SX-ACE Processor

Ryusuke Egawa (Tohoku University), Shintaro Momose (NEC Corporation), Kazuhiko Komatsu (Tohoku University), Yoko Isobe (NEC Corporation), Hiroyuki Takizawa (Tohoku University), Akihiro Musa (NEC Corporation), Hiroaki Kobayashi (Tohoku University)

Using practical scientific and engineering applications, this poster presents early performance evaluation of the SX-ACE vector processor, which is the latest vector processor developed by NEC in 2013. While inheriting the advantages of the vector architecture, the SX-ACE processor is designed so as to overcome the drawbacks of conventional vector processors by introducing several architectural features. To unveil the potential of the SX-ACE processor, this poster evaluates and discusses how the brand-new vector processor is beneficial to achieve a high sustained performance on practical science and engineering simulations. Evaluation results show that, for a wide range of practical simulations, SX-ACE potentially achieves a higher sustained performance than its predecessor SX-9 and existing scalar processors.

Tightly Coupled Accelerators Architecture for Low-latency Inter-Node Communication Between Accelerators

Toshihiro Hanawa (University of Tokyo), Yuetsu Kodama, Taisuke Boku, Mitsuhisa Sato (University of Tsukuba)

Inter-node communications between accelerators in heterogeneous clusters require extra latency due to the data copy between host and accelerator, and such communication latency causes severe performance degradation on applications. Espe-

cially in the next generation's HPC systems, the strong scaling will be more serious issue than today, and the communication latency becomes the critical issue. To address this problem, we proposed the Tightly Coupled Accelerators (TCA) architecture, and designed the interconnection router chip named PEACH2. Accelerators in the TCA architecture communicate directly via the PCIe protocol to eliminate protocol overhead, as well as the data copy overhead. In this paper, we present HA-PACS/TCA system, the proof-of-concept GPU cluster based on the TCA architecture. Our system demonstrates 2.3 μ sec of the latency on the inter-node GPU-to-GPU communication. As the result of Himeno benchmark, we demonstrated that TCA improves the scalability of the performance in the small size with up to 65%.

Award: Best Poster Finalist

Employing Machine Learning for the Selection of Robust Algorithms for the Dynamic Scheduling of Scientific Applications

Nitin Sukhija, Srishti Srivastava (Mississippi State University), Florina M. Ciorba (Technical University Dresden), Ioana Bancicescu (Mississippi State University), Brandon Malone (Helsinki Institute for Information Technology)

Scheduling scientific applications with large, computationally intensive, and data parallel loops, which have irregular iteration execution times, on heterogeneous computing systems with unpredictably fluctuating load requires highly efficient and robust scheduling algorithms. State-of-the-art dynamic loop scheduling (DLS) techniques provide a solution for achieving the best performance for these applications executing in dynamic computing environments. Selecting the most robust of the state-of-the-art DLS algorithms remains, however, challenging.

In this work we propose a methodology for solving this selection problem. We employ machine learning to obtain an empirical robustness prediction model that enables algorithm selection from a portfolio of DLS algorithms on a per-instance basis. An instance consists of the given application and current system characteristics, including workload conditions. Through discrete event simulations, we show that the proposed portfolio-based approach offers higher performance guarantees with respect to the robust execution of the application when compared to the simpler winner-take-all approach.

Large-Scale Parallel Visualization of Particle Datasets Using Point Sprites

Silvio Rizzi, Mark Hereld, Joseph Insley, Michael E. Papka, Thomas Uram, Venkatram Vishwanath (Argonne National Laboratory)

Massive particle-based datasets can be efficiently rendered using point sprites. Among other advantages, the method is efficient in its use of memory, preserves the dynamic range of the

simulations, and provides excellent image quality. Preliminary results of large-scale parallel rendering of particle data sets using point sprites are presented in this poster. Performance and scalability are evaluated on a GPU-based computer cluster using datasets of up to 3200^3 particles at resolutions of up to 6144×3072 pixels.

A Multiple Time Stepping Algorithm for Efficient Multiscale Modeling of Platelets Flowing in Blood Plasma

Na Zhang, Yuefan Deng (Stony Brook University)

Numerical simulation of complex biological systems is an immense computational and algorithmic challenge primarily due to modeling complexity and disparate spatiotemporal scales of the mechanisms occur. We develop a parallel multiscale multiple time-stepping (MTS) algorithm for exploring the mechanisms at disparate spatiotemporal scales. Specifically, we apply our algorithm on parallel computers with tens of thousands of cores to study flow-induced platelet-mediated thrombogenicity. This MTS algorithm improves considerably computational efficiency without significant loss of accuracy. This poster presents mathematical models and multiscale methods for study of dynamic properties of flowing platelets, multiscale parallel algorithms for reducing computing time, implementation, and quantitative analysis of performance metrics versus accuracy metrics. This MTS algorithm establishes a computationally feasible approach for performing efficient multiscale simulations.

PyFR: An Open Source Python Framework for High-Order CFD on Heterogeneous Platforms

Freddie D. Witherden, Brian C. Vermeire, Peter E. Vincent (Imperial College London)

Computational fluid dynamics underpins several high-tech industries. However, the numerics utilized by most industrial codes are not suitable candidates for acceleration. In this work we present a framework for solving the compressible Navier-Stokes equations efficiently across a range of hardware platforms. The framework, PyFR, is open source and written entirely in Python. By employing the flux reconstruction approach of Huynh PyFR is able to combine the accuracy of spectral schemes with the geometric flexibility of low-order finite volume schemes.

Our results show that PyFR is both accurate and performance portable across CPUs and GPUs from Intel, AMD and NVIDIA. We also demonstrate the heterogeneous capabilities of PyFR which permit a simulation to be decomposed across a range of platforms. Finally, we show the scalability of PyFR on up to 104 NVIDIA M2090 GPUs.

The Parallel Java 2 Library: Parallel Programming in 100% Java

Alan R. Kaminsky (Rochester Institute of Technology)

The Parallel Java 2 Library (PJ2) is an API and middleware for parallel programming in 100% Java on multicore parallel computers, cluster parallel computers, and hybrid multicore cluster parallel computers. In addition, PJ2 supports programming GPU accelerated parallel computers, with main programs written in Java and GPU kernels written using NVIDIA's CUDA. PJ2 also includes a lightweight map-reduce framework for big data parallel programming. PJ2 has its own job scheduling middleware that coordinates usage of a node's or cluster's computational resources by multiple users. Encompassing four major categories of parallel programming—multicore, cluster, GPU, and map-reduce—in a single unified Java-based API, PJ2 is suitable for teaching and learning parallel programming and for real world parallel program development.

Performance Grading of GPU-Based Implementation of Space Computing Systems Image Compression

Olympia Kremmyda (National and Kapodistrian University of Athens), Vasilis Dimitsas (National and Kapodistrian University of Athens), Dimitris Gizopoulos (National and Kapodistrian University of Athens)

Modern GPUs can accelerate the execution of inherently data-parallel applications. However, real-world algorithm realizations on GPUs differ from synthetic benchmarks. In this paper we stress the performance limits of GPUs. We measure the performance of our CUDA implementation of a recommended image compression algorithm (CCSDS-122.0-B-1) for space data systems on NVIDIA GPUs for various image sizes and compare it to the execution on x86 CPUs. The control-intensive parts and irregular memory access patterns under-utilize the GPU architecture and effectively serialize large parts of the GPU kernel execution. We present the execution results on two systems with different CPU/GPU setups. In both systems, our GPU implementation of the algorithm eventually leads to a very moderate 1.1x to 2.1x speedup. Discussion of the performance bottlenecks leads to insights on how GPU execution can be improved in the future by either revising the algorithms or by supporting control-intensive execution on the GPU cores.

Performance Portable Parallel Programming - Compile-Time Defined Parallelization and Storage Order for Accelerators and CPUs

Michel Müller (Tokyo Institute of Technology)

Performance portability between CPU and accelerators is a major challenge for coarse grain parallelized codes. Hybrid Fortran offers a new approach in porting for accelerators that requires minimal code changes and keeps the performance of CPU optimized loop structures and storage orders. This is achieved through a compile-time code transformation where the CPU

and accelerator cases are treated separately. Results show minimal performance losses compared to the fastest non-portable solution on both CPU and GPU. Using this approach, five applications have been ported to accelerators, showing minimal or no slowdown on CPU while enabling high speedups on GPU.

Cosmography with SDvision

Daniel Pomarède (CEA Saclay), Hélène Courtois (University Claude Bernard I), Yehuda Hoffman (Hebrew University of Jerusalem), R. Brent Tully (University of Hawaii, Honolulu)

Cosmography is the creation of maps of the Universe. Using the SDvision 3D visualization software developed within the framework of IDL Object Graphics, we have established a cosmography of the Local Universe, based on multiple data products from the Cosmic Flows Project. These data include catalogs of redshifts, catalogs of peculiar velocities, and reconstructed density and velocity fields. On the basis of the various visualization techniques offered by the SDvision software, that rely on multicore computing and OpenGL hardware acceleration, we have created maps displaying the structure of the Local Universe where the most prominent features such as voids, clusters of galaxies, filaments and walls, are identified and named. Maps also display the dynamical information of the cosmic flows, which are the bulk motions of galaxies, of gravitational origin. These maps highlight peculiar conformations in the cosmic flows such as the streaming along filaments, or the existence of local attractors.

Machine Learning Algorithms for the Performance and Energy-Aware Characterization of Linear Algebra Kernels on Multithreaded Architectures

A. Cristiano I. Malossi, Yves Ineichen, Costas Bekas, Alessandro Curioni (IBM Zurich Research Laboratory), Enrique S. Quintana-Ortí (James I University)

The performance and energy optimization of the 13 “dwarfs”, proposed by UC-Berkeley in 2006, can have a tremendous impact on a vast number of scientific applications and existing computational libraries. To achieve this goal, scientists and software engineers need tools for analyzing and modeling the performance-power-energy interactions of their kernels on real HPC systems.

In this poster we present systematic methods to derive reliable time-power-energy models for dense and sparse linear algebra operations. Our strategy is based on decomposing the kernels into sub-components (e.g., arithmetics and memory accesses) and identifying the critical features that drive their performance, power, and energy consumption. The proposed techniques provide tools for analyzing and reengineering algorithms for the desired power- and energy-efficiency as well as to reduce operational costs of HPC-supercomputers and cloud-systems with thousands of concurrent users.

Visualization of Particles Beam Simulations in the IFMIF Accelerator

Bruno Thooris, Phu-Anh-Phi Nghiem, Daniel Pomarède (CEA Saclay, Institute of Research into the Fundamental Laws of the Universe)

Using the SDvision 3D visualization software developed within the framework of IDL Object Graphics at CEA/IRFU, we have visualized the travel of particles in the IFMIF-LIPAc accelerator. IFMIF (Internal Fusion Materials Irradiation Facility) is a Europe-Japan joint project aiming at constructing an accelerator-based neutron source, the world’s most intense one, dedicated to study materials that must withstand the intense neutron flux coming from the fusion plasma of future Tokamaks. Tokamaks are the nuclear fusion reactors capable of producing energy in a similar way as in the Sun’s core. Simulations visualized here were made with a million macroparticles. To ensure that beam losses do not exceed 1W/m in the high energy part, simulations with 4,000,000,000 particles were performed, which lasted several weeks using fifty processors in parallel. Stereoscopic video sequences in 3D were achieved at IRFU and a 15-minute movie has been realized for outreach purposes.

Performance Model for Large-Scale Simulations with NEST

Wolfram Schenck, Andrew V. Adinets, Yuri V. Zaytsev, Dirk Pleiter, Abigail Morrison (Juelich Supercomputing Centre)

NEST is a simulator for large networks of spiking point neurons for neuroscience research. A typical NEST simulation consists of two stages: first the network is wired up, and second the dynamics of the network is simulated. Our work aims at developing a performance model for the second stage, the simulation stage, by a semi-empirical approach. We collected measurements of the runtime performance of NEST under varying parameter settings on the JUQUEEN supercomputer at Forschungszentrum Jülich, and subsequently fitted a theoretical model to this data. This performance model defines the simulation time as weighted sum of algorithmic complexities which have been identified in the NEST source code. After parameter fitting, the coefficient of determination on the training data is close to 1.0, and the model can be used to successfully extrapolate NEST simulation times. Furthermore, recommendations for algorithmic improvements of the NEST code can be derived from the modeling results.

Interoperating MPI and Charm++ for Productivity and Performance

Nikhil Jain (University of Illinois at Urbana-Champaign), Abhinav Bhatele, Jae-Seung Yeom (Lawrence Livermore National Laboratory), Mark F. Adams (Lawrence Berkeley National Laboratory), Francesco Miniati (ETH Zurich), Chao Mei (Google), Laxmikant Kale (University of Illinois at Urbana-Champaign)

This poster studies interoperation among two parallel languages that differ with respect to the driver of program execution, Charm++ and MPI. We describe the challenges in enabling interoperation between MPI, a user-driven language, and Charm++, a system-driven language. We present techniques for managing important attributes of a program, such as the control flow, resource sharing, and data sharing, in an interoperable environment. We show that our implementation enables interoperation between MPI and Charm++ via minor addition to the source code and promotes reuse of existing software. Finally, we study the application of the presented techniques and demonstrate the benefits of interoperation through several case studies using production codes including Chombo, EpiSimdemics, NAMD, FFTW, and MPI-IO, executed on IBM Blue Gene/Q and Cray XE6.

Performance Optimization and Evaluation of a Global Climate Application Using a 440m Horizontal Mesh on the K Computer

Masaaki Terai, Hisashi Yashiro (RIKEN), Kiyotaka Sakamoto (Fujitsu Systems East Limited), Shin-ichi Iga (Computational Climate Science Research Team, Advanced Institute for Computational Science, RIKEN), Hirofumi Tomita (RIKEN), Masaki Satoh (University of Tokyo), Kazuo Minami (RIKEN)

Using almost all the compute nodes on the K computer, we evaluated the performance of NICAM, a climate application for a high-resolution global atmospheric simulation study. First, based on profiling results, we improved single node performance using kernel programs extracted from the application, resulting in increases in the peak performance ratio of the kernel programs. Thereafter, we applied the optimization results to the application, and evaluated the performance by weak scaling. In the evaluation, to retain the same problem size per node, we reduced the horizontal mesh size as the number of nodes increased, decreasing it to 0.44km with 81,920 nodes. Although we achieved good scalability of the dynamical step, the elapsed time of the physical step increased with the number of nodes owing to inter-node load imbalances in the cloud microphysics. Finally, the peak performance ratio of the main loop of NICAM reached 8.3% with 81,920 nodes.

A Roofline Performance Analysis of an Algebraic Multigrid Solver

Alex Druinsky, Brian Austin, Xiaoye S. Li, Osni Marques, Eric Roman, Samuel Williams (Lawrence Berkeley National Laboratory)

We present a performance analysis of a novel element-based algebraic multigrid (AMGe) method combined with a robust coarse-grid solution technique based on HSS low-rank sparse factorization. Our test datasets come from the SPE Comparative Solution Project for oil reservoir simulations. The current performance study focuses on one multicore node and on bound analysis using the roofline technique. We found that keeping a small value of spectral tolerance is most critical to achieve the best AMG solver performance. Our roofline bound estimate is within 23% accuracy compared to the actual runtime on a Cray XC30 12-cores processor.

bgclang: Creating an Alternative, Customizable, Toolchain for the Blue Gene/Q

Hal Finkel (Argonne National Laboratory)

bgclang, a compiler toolchain based on the LLVM/Clang compiler infrastructure, but customized for the IBM Blue Gene/Q (BG/Q) supercomputer, is a successful experiment in creating an alternative, high-quality compiler toolchain for non-commodity HPC hardware. By enhancing LLVM with support for the BG/Q's QPX vector instruction set, bgclang inherits from LLVM/Clang a high-quality auto-vectorizing optimizer, C++11 frontend, and many other associated tools. Moreover, bgclang provides both unique capabilities and a much-needed platform for experimentation. Here we'll focus on two examples: First, bgclang's Address Sanitizer feature, which provides efficient instrumentation-based memory-access validation at scale; Second, experiments in loop optimization techniques, including a software prefetching optimization, within bgclang. These loop optimizations can produce 2x speedups over code produced by other available compilers for loops in the TSVC benchmark suite.

New Parallelization Model of Sequential Monte Carlo Analysis with Prediction-Correction Computing

Eiji Tomiyama, Hiroshi Koyama (Research Organization for Information Science and Technology), Katsumi Hagita (National Defense Academy)

Monte Carlo (MC) search with Metropolis judgment is widely used to solve inverse problems. In particular Reverse Monte Carlo (RMC) analysis reproduces particle's configuration in terms of the structure factor $S(q)$ obtained by x-ray scattering experiments. In general, MC search proceeds sequentially because each judgment of the MC trial depends on its previous state. In this study, we have developed a new 'SimpleRMC'

parallel code for RMC analysis with millions of particle system. Two hotspots are identified and optimized at the histogram $h(r)$ and difference of histogram $\Delta h(r)$ calculations. In the histogram kernel, we achieved 339 Tflops performance (32.3 % of the peak) on 8,192 nodes of K computer using 33,554,432 particle system. The 'prediction-correction' method improves parallel performance of $\Delta h(r)$ calculation as well as its elapse time.

Automata Processor Accelerates Association Rule Mining

Ke Wang, Kevin Skadron (University of Virginia)

Association rule mining (ARM) is a widely used data mining technique. As the databases grow quickly, the performance becomes a critical issue for ARM. We present an acceleration solution to speedup association rule mining problem by using Micron's Automata Processor (AP) mapping itemset recognition to a non-deterministic finite automaton. High parallelism is achieved by utilizing a huge amount of pattern matching and counting elements on the AP board. Apriori framework is adopted to prune the search space of itemset candidates. Several performance optimization strategies are proposed to improve the performance including data pre-processing, AP I/O optimization and trading slow routing recompilation with fast symbol-replacement. 10X-20X speedup is achieved for real benchmark databases when compared with the Apriori CPU implementation.

Leveraging Naturally Distributed Data Redundancy to Optimize Collective Replication

Bogdan Nicolae, Massimiliano Meneghin, Pierre Lemarinier (IBM Research)

Techniques such as replication are often used to enable resilience and high availability. However, replication introduces overhead both in terms of network traffic necessary to distribute replicas, as well as extra storage space requirements. To this end, redundancy elimination techniques such as compression or deduplication are often used to reduce the overhead of communication and storage. This paper aims to explore how these two phases can be optimized by combining them into a single phase. Our key idea relies on the observation that since data is related, there is a probability that distributed redundancy is already naturally present, thus it may pay off to try to identify this natural redundancy in order to avoid reducing redundancy unnecessarily in the first phase only to add it back later in the second phase. We present how this idea can be leveraged in practice and demonstrate its viability for two real-life HPC applications.

RHF SCF Parallelization in GAMESS by Using MPI and OpenMP

Yuri Alexeev, Graham Fletcher, Vitali Morozov (Argonne National Laboratory)

In the exascale era, only carefully tuned applications are expected to scale to millions of cores. For quantum chemistry, one of the practical solutions to achieve such scalability is to employ new highly parallel methods which use both MPI and OpenMP. In this work, we optimized code and developed a thread safe Rys quadrature integral code with parallelized Self Consistent Field method using OpenMP in quantum chemistry package GAMESS. We implemented and benchmarked three different hybrid MPI/OpenMP algorithms, and discussed their advantages and disadvantages. The most efficient MPI/OpenMP algorithm was benchmarked against an MPI-only implementation on the Blue Gene/Q supercomputer and was shown to be consistently faster by a factor of two.

Space-Filling Curves for Domain Decomposition in Scientific Simulations

Aparna Sasidharan, Marc Snir (University of Illinois at Urbana-Champaign)

In this work we explore the possibility of using space-filling curves (SFC) as a quick and easy method to produce good quality mesh partitions. The existing algorithms for generating SFCs are limited by the geometry and size of the domain. We propose a recursive algorithm that can be used to generate a general SFC for domains of arbitrary shapes and sizes. We provide rules for generating SFCs for 2D domains and extend them to support domains in higher dimensions. The meshes we used as test cases come from Community Earth System Model (CESM). The quality of partitions is estimated based on their maximum computation and communication loads. We compared the SFC partitions with the existing strategies in CESM, including the multi-level k-way partitioning algorithms of Metis.

Lossy Compression for Checkpointing: Fallible or Feasible?

Xiang Ni (University of Illinois at Urbana-Champaign), Tanzima Islam, Kathryn Mohror (Lawrence Livermore National Laboratory), Adam Moody, Laxmikant Kale (University of Illinois at Urbana-Champaign)

As HPC applications scale to hundreds of thousands of processors, large checkpoints consume a lot of space making it costly to fit them in memory or burst buffers. It also takes a significant amount of time to transfer them to stable storage. Lossless compression fails to reduce the size of such checkpoints due to randomness in the lower bits of typical floating point scientific data. To address this challenge, we propose use

of lossy compression for reducing checkpoint size significantly, and study its impact on correctness of application execution. We study the trade-off between the loss of precision, compression ratio, and correctness due to lossy compression. For ChaNGa, we show that use of moderate lossy compression reduces checkpoint size by 3-5x while maintaining correctness. Finally, we inject failures following different distributions to study whether an application is more sensitive to precision loss at earlier or later stages of the simulation.

Toward Effective Detection of Silent Data Corruptions for HPC Applications

Sheng Di (Argonne National Laboratory), Eduardo Berrocal (Illinois Institute of Technology), Leonardo Bautista-Gomez, Katherine Heisey, Rinku Gupta, Franck Cappello (Argonne National Laboratory)

Because of the large number of components, future extreme-scale systems are expected to suffer a lot of silent data corruptions. Changes caused by silent errors flipping low-order bit positions are very small, making them difficult to detect by software. In this work, we convert the detection problem to a one-step look-ahead prediction issue and explore the most effective prediction methods for different HPC applications. We exploit the Auto Regressive (AR) model, Auto Regressive Moving Average (ARMA) Model, Linear Curve Fitting (LCF), and Quadratic Curve Fitting (QCF). We evaluate them using real HPC application traces. Experiments show that the error feedback control plays an important role in improving detection. AR and QCF perform the best among all evaluated methods, where F-measure can be kept around 80% for silent bit-flip errors occurring around the bit position 20 for double-precision data or around bit 8 for single-precision data.

Hardware Accelerated Linear Programming: Parallelizing the Simplex Method with OpenCL

Bradley de Vlugt (Western University), Maysam Mirahmadi (IBM Canada Ltd.), Serguei L. Primak (Western University), Abdallah Shami (Western University)

This work proposes an energy-efficient hardware accelerated Linear Programming (LP) solver. The system is based on the Simplex Algorithm for solving LP problems and is operable on Field Programmable Gate Arrays (FPGAs), Graphic Processing Units (GPUs), and Multi-Core Computer Processors (CPUs). The system is targeted towards the dense problems in radio-therapy applications as they represent a challenge to modern solvers.

Performance benchmarking reveals speed ups relative to a sequential implementation that approach 2 and 10 on a CPU and GPU for random, dense problems. The FPGA exhibits unity speed up but proved to be the most efficient in terms of

Simplex iterations processed per unit energy with efficiency 5 times greater than the CPU. This is a notable speed improvement and power saving in comparison with current technology for solving dense problems as the GPU code can solve problems with speeds up to 50 times greater than a sparse solver.

Power Shifting Opportunities on BG/Q Using Memory Throttling

Bo Li (Virginia Polytechnic Institute and State University), Edgar Leon (Lawrence Livermore National Laboratory)

Power and energy efficiency are major concerns for future exascale systems. Achieving the expected levels of application performance in these over-provisioned, power-constrained systems will be challenging. Thus, understanding how power is consumed by an application throughout its different phases will be necessary to shift power to those resources on the critical path and maximize performance under the operating power budget. In this project, we identify power shifting opportunities for two proxy-applications at Lawrence Livermore National Laboratory: LULESH and the kernels of pF3D. With a hardware-counter based linear regression model, we dynamically throttle the memory system on a per-region or per-kernel basis on a Blue Gene/Q system. Our results show that we can save up to 20% of power on regions with low memory bandwidth requirements at a marginal performance cost. Critical path resources could use this power to improve application performance.

ACM Student Research Competition Posters

WrAP: Write Aside Persistence for Storage Class Memory in High Performance Computing

Ellis Giles (Rice University)

Emerging memory technologies like Phase Change Memory or Memristors (generically called SCM or Storage Class Memory) combine the ability to access data at byte granularity with the persistence of storage devices like hard disks or SSDs. By accessing data directly from SCM addresses instead of slow block I/O operations, developers can gain 1-2 orders of performance. However, this pushes upon developers the burden of ensuring that SCM stores are ordered correctly, flushed from processor caches, and if interrupted by sudden machine stoppage, do not leave objects in SCM in inconsistent states.

Software based Write Aside Persistence, or WrAP, provides durability and consistency for SCM writes, while ensuring fast paths to data in caches, DRAM, and persistent memory tiers. Simulations such as the Graph 500 Benchmark indicate significant performance gains over traditional log based approaches.

Archer: A Low Overhead Data Race Detector for OpenMP

Simone Atzeni (University of Utah)

On extreme-scale systems, large applications are increasingly turning to OpenMP to harness an explosion of on-node parallelism. However, code complexity in large applications can lead programmers to use OpenMP incorrectly, and data races are particularly challenging to diagnose. Current approaches to identify races include dynamic and static analyses. Unfortunately, either technique alone is often ill suited for large applications as it suffers from high run-time overhead and in-accuracy. We present Archer, a novel approach that combines static and dynamic techniques to identify data races in large OpenMP applications with low overhead while still providing high accuracy. Our LLVM-based static analysis techniques classify all of the OpenMP regions within an application and pass only those potentially unsafe regions to our run-time component that extends Google's ThreadSanitizer. Our preliminary results show that reduction in performance and memory overhead of our run-time component is proportional to the amount of code it must examine.

Analysis of Energy and Power Consumption by Remote Data Transfers in Distributed Programming Models

Siddhartha Jana (University of Houston)

Recent studies of the challenges facing the Exremescale era express a need for understanding the impact on energy profile of applications due to inter-process communication on large-scale systems. Programming models like MPI provide the user with explicit interfaces to initiate data-transfers among distributed processes.

This poster presents empirical evidence that: controllable factors like the size of the data-payload to be transferred and the number of explicit calls used to service such transfers have a direct impact on the power signatures of communication kernels. Moreover, the choice of the transport layer (along with the associated interconnect) and the design of the inter-process communication protocol are responsible for these signatures. Results discussed in this poster motivate the incorporation of energy-based metrics for fine tuning communication libraries that target Exremescale machines.

Computing Approximate b-Matchings in Large Graphs and an Application to k-Anonymity

Arif Khan (Purdue University)

Given a graph, b-Matching problem is to an edge weighted matching with maximum weight with constraint that every vertex v can match with at most a specified number of b vertices. It has been shown that b-Matching is useful in various machine learning problems such as classification, spectral clustering, graph sparsification, graph embedding and data privacy. The exact algorithms for this problem have high time as well as storage requirements, are inherently sequential, and therefore, are not practical solving large problems. We propose a $1/2$ -approximation algorithm which runs in linear time in the number of edges and also requires less storage. We show that our algorithm can solve large problems and can get up to 96% of the optimal solution. We also show that our algorithm scales up to 10x on 16 cores of Intel Xeon machines and up to 50x on 60 cores of Intel Xeon Phi machines.

A Framework for Distributed Tensor Computations

Martin D. Schatz (University of Texas at Austin)

Recently, data models have become more complex leading to the need for multi-dimensional representations to express the data in a more meaningful way. Commonly, tensors are used to represent such data and multi-linear algebra, the math associated with tensors, has become essential for tackling problems in big data and scientific computing. Up to now, the main approach to solving problems of multi-linear algebra has been based on mapping the multi-linear algebra to linear algebra and relying on highly efficient linear algebra libraries to perform the equivalent computation. Unfortunately, there are inherent inefficiencies associated with this approach. In this work, we define a notation for tensor computations performed on distributed-memory architectures. Additionally, we show how, using the notation, algorithms can be systematically derived, required collective communications identified, and approximate costs analyzed for a given tensor contraction operation.

Supervised Learning for Parallel Application Performance Prediction

Andrew Titus (Massachusetts Institute of Technology)

Communication is a scaling bottleneck for many parallel applications executed on large machines. Intelligent task mapping may alleviate the negative impact of communication, but simple metrics used to find such mappings may not be good predictors of their performance. We evaluate supervised machine learning methods as tools for prediction of communication time of large parallel applications. Through these methods, we correlate communication time for different task

mappings to the corresponding network hardware counters accessible on the IBM Blue Gene/Q. The results from these machine learning regression algorithms are used to provide insight into the relative importance of different hardware counters and metrics for predicting application performance. These results are explored graphically in the poster in the context of two production applications, MILC and pF3D.

Performance Variability Due to Job Placement on Edison

Dylan Wang (University of California, Davis)

Some applications running on machines like the Edison Supercomputer can suffer from high variability in run-time. This leads to debugging and optimization difficulties and less accurate reservation times. Supercomputers with high performance variability end up being less useful for end users and less efficient for the supercomputing facility. The objective of this research is to characterize the application run-time performance and identify the root cause of the variability on machines with the Aries network. We approach this problem by running two applications, MILC and AMG, while collecting the set of logical coordinates for every job's nodes in the system and hardware counters on routers connected to our application's nodes. Running these applications at various sizes once per day will give us many unique allocations and system states with corresponding execution times. We then use statistical analysis on the gathered data to try and correlate performance with placement and interference.

Multi-Level Hashing Dedup in HPC Storage Systems

Eric E. Valenzuela (Texas Tech University)

Reaching high ratios of data deduplication in High Performance Computing (HPC) is highly achievable. Prior art demonstrates magnitudes of reduction possible and 15 to 30 percent of redundant data can be removed on average using deduplication techniques. The objective of this research study is to design and experiment a dedup system to provide 100% data integrity without a possibility of losing data while reducing the need of costly byte-by-byte comparisons. Because data deduplication uses hashing algorithms, hash collisions will occur. Prior systems ignore byte-by-byte comparisons that are needed to handle collisions citing the probability is low. Our research focuses on investigating a multi-level dedup method to reduce byte-by-byte comparisons while providing 100% data integrity, and the implementation of multi-level hash functions while taking advantage of Xeon Phi many-core architecture to compute cryptographic fingerprints concurrently. Our current proof-of-concept evaluations with a deduplication file system, Lessfs, show promising results.

Parallel Complex Network Partitioning

George M. Slota (Pennsylvania State University)

A large number parallel graph analytics follow a bulk synchronous parallel (BSP) model: periods of parallel computation followed by periods of parallel communication. To maximize parallel efficiency when a graph is distributed across a cluster, we want a partitioning of the input graph that balances both work and memory (proportional to number of vertices and edges per process) and communication (proportional to total edge cut and maximal edge cut per process). Traditional multi-level partitioners are unable to satisfy all of these requirements and are heavy-weight in terms of computational and memory requirements. This work introduces PULP, an iterative partitioning methodology for small-world graphs that can simultaneously handle multiple constraints and multiple objectives. Partitions produced by PULP are equal to or better than state-of-the-art partitioners in terms of edge cut. PULP also runs very fast, being capable of partitioning multi-billion edge graphs in minutes on a single compute node.

Enhancing FlashSim Simulations for High Performance Computing Storage Systems

Omar Rodriguez (Polytechnic University of Puerto Rico)

This work describes the design and development of an enhanced FlashSim simulator framework and toolkit. The FlashSim is one of the most popular SSD simulators. SSD simulator focuses on software components, specifically FTL (flash translation layer) schemes, garbage collection, and wear-leveling policy. In order to simulate and analyze the performance results, the simulator needs block I/O traces. This work extends the FlashSim simulator with leveraging the blktrace tool to automate collecting traces from different application workloads. We also develop a converter that transforms the I/O traces into the format needed for the FlashSim simulator. The proposed extension enhances the FlashSim simulator and automates trace collection, conversion, and analysis.

Orthogonal Scheduling of Stencil Computations with Chapel Iterators

Ian Bertolacci (Colorado State University)

Stencil computations are an important element of scientific computing that can suffer from performance caps due to poor iteration scheduling. Implementing stencil computations is already a difficult task, and the addition of specifying schedules with both good parallelism and high data locality introduces a new dimension of difficulty and code complexity. While compilers already exist to automate the task of optimization, they may lack the information required to recognize when code could be transformed to take advantage of a particular scheduling.

Using the Chapel language, we implemented tiling schedules decoupled from the implementation of the stencil computation. We observed performance gains equivalent to highly coupled code, without introducing the complexity. We will discuss the advantages of using iterators as a scheduling mechanism, as well as the possibility of creating a Chapel module of iterators that would provide developers with a simple method to include better optimized scheduling into their code.

Floating-Point Robustness Estimation by Concrete Testing

Wei-Fan Chiang (University of Utah)

Analyzing floating-point imprecision is required for dealing with the performance-precision trade-off of numerical programming. However, without specific mathematical knowledge, this is hard to be done by general developers. We focus on solving one of the many floating-point precision issues, robustness estimation, by concrete testing. Programs without sufficient robustness tend to generate unexpected outputs, but developers cannot tune their programs without knowing the robustness degrees of candidate routines. Our testing driven framework estimates robustness by enumerating the inputs causing non-robust behavior. The number of problematic inputs indicates robustness degree and provides a comparison basis between candidate routines. We evaluated our framework with two implementations of a geometric primitive: one of the two implementations is pointed out to be more robust than another. This empirical result matches the robustness argument in the previous context. Thus we strongly believe that it is worthy to further investigate concrete testing based robustness estimation.

Gklepp: Parallelizing a Symbolic GPU Race Checker

Mark S. Baranowski (University of Utah)

The increased usage of GPGPU programs necessitates tools to detect errors in the program. Existing tools which use hardware instrumentation may miss certain bugs due to limitations in the memory they observe and the requirement that tests must produce the bug. Formal methods found in tools such as Gklepp can find data races without users providing test cases. The primary disadvantage of using such tools is the slowness of execution when compared to the speed at which conventional tools run. We present a parallelized version of Gklepp, christened Gklepp, which better uses computation resources through parallel execution. We can attain speedups of 2-8 times while retaining most detection of data races.

Reducing Network Contention Associated with Parallel Algebraic Multigrid

Amanda J. Bienz (University of Illinois at Urbana-Champaign)

Algebraic multigrid (AMG) is an iterative method often used to solve PDEs arising in various fields of science and engineering. The method is optimal, requiring only $O(n)$ work to solve a system of n unknowns. Parallel implementations of AMG, however, lack scalability. As problem size increases, parallel AMG suffers from increasingly dense communication patterns, yielding network contention and reduced efficiency. The amount of communication can be reduced algorithmically with little change in convergence, through use of non-Galerkin coarse grids. The overall performance of parallel AMG can be further improved upon through the tradeoff of communication and convergence rates. Removing costly communication, such as that between two processors on opposite ends of the network, can improve the runtime of AMG if the cost saved from communication is greater than that added from decreasing convergence.

A Molecular Model for Platelets at Multiple Scales and Simulations on Supercomputers

Li Zhang (Stony Brook University)

A molecular model at multiscale for human platelets is designed and implemented on large-scale supercomputers. The model characterizes the principal physiological functions of key components of single platelet. The organelle zone centralizes as the cytoplasm, the peripheral zone includes the bilayer and exterior coats, and the cytoskeleton zone forms the flow-induced filopodia. It is the first time such a massive model in biomedical engineering is proposed and high-performance computing is the only potential solution for initial-understanding of the model. Algorithmically, a coarse-grained molecular dynamics force field describes the molecular-level interactions in the membrane and cytoskeleton while a Morse potential is applied to the cytoplasm. Together, a coarse-grained stochastic dynamics is applied to simulating the macroscopic viscous fluid flows. Numerical experiments with our model require considerations of large number of particles and the physiological phenomena at multiscale and thus demand development of the-state-of-the-art parallel algorithms, the main thrust of this research.

Adaptive Power Efficiency: Runtime System Approach with Hardware Support

Ehsan Toton (University of Illinois)

Power efficiency is an important challenge for the HPC community and demands major innovations. We use extensive application-centric analysis of different architectures to design automatic adaptive runtime system (RTS) techniques that save significant power. These techniques exploit common application patterns and only need minor hardware support. The application's pattern is recognized using formal language theory to predict its future and adapt the hardware appropriately.

I will discuss why some system components such as caches and network links consume extensive power disproportionately for common HPC applications, and how a large fraction of power consumed in caches and networks can be saved using our approach automatically. In these cases, the hardware support the RTS needs is the ability to turn off ways of set-associative caches and network links. I will also give an overview of an RTS approach for handling process variation.

Towards Improving the Performance of the ADCIRC Storm Surge Modeling Software

Nick Weidner (University of South Carolina)

The objective of this work was to run the ADVanced CIRCulation (ADCIRC) storm surge modeling software on the Stampede Super Computer at the Texas Advanced Computing Center. In order to take advantage of the Intel Xeon Phi Co-Processors, we identified code hotspots from the host processors and then increased the speed by implementing OpenMP on parallel work. We then worked to offload the threaded work on the Xeon-Phi in order to identify speed up in comparison to previous test.

Creating a Framework for Systematic Benchmarking of High Performance Computing Systems

Mike Pozulp (College of William & Mary)

The Applications Performance & Productivity (APP) Group of the High-End Computing Capability (HECC) Project is tasked with ensuring maximal performance of the supercomputing systems operated at NASA Ames. As part of this task, one focus is to benchmark current and future systems so as to guide users to the most suitable resources for their applications. The challenges of benchmarking include: intractable combinations of parameters and system configurations, issues in organization and persistence of benchmarking results, and problems of determining benchmark performance or computational success. The benchmarking workflow has three essential steps: parameter selection, code execution, and result interpretation. A novel implementation of a semi-automated test framework

has been utilized to analyze performance variation of the NAS Parallel Benchmark (NPB) applications running on the resident supercomputer, Pleiades. The framework has also been used to test the hypothesis that reductions in performance are correlated with the corresponding variation in network communication distance.

Scalable Fault Tolerance in Multiprocessor Systems

Gagan Gupta (University of Wisconsin-Madison)

Evolving trends in design and use of computers are resulting in fault-prone systems which may not execute a program to completion. Checkpoint-and-recovery (CPR) is commonly used to recover from faults and complete parallel programs. Conventional CPR incurs high overheads and may be inadequate in the future as faults become frequent. This work proposes to execute parallel programs deterministically to enable lower overhead and scalable fault tolerance.

Static Analysis of MPI Programs Targeting Parallel Properties

Sriram Ananthakrishnan (University of Utah)

MPI is one of the dominant programming model for writing HPC applications. Unfortunately, debugging MPI programs is hard. It is well known that static dataflow analysis discovers provably true properties that are useful in optimization and proving correctness. Existing techniques fail to interpret MPI operations for their message passing semantics thereby lacking the ability to recognize parallel properties of MPI programs. Parallel properties are properties about communication topology or properties that depend on communication; e.g., reaching constants over MPI operations. In this work we propose (i) new abstractions and techniques to statically discover the communication topology (ii) provide a framework for writing dataflow analyses to recognize parallel properties of MPI programs.

Real-Time Outlier Detection Algorithm for Finding Blob-Filaments in Plasma

Lingfei Wu (College of William & Mary)

Magnetic fusion could be an inexhaustible, clean, and safe solution to the global energy needs. The success of magnetically confined fusion reactors demand steady-state plasma confinement which is challenged by the edge turbulence such as the blob-filaments. Real time analysis can be used to monitor the progress of fusion experiments and prevent catastrophic events. We present a real-time outlier detection algorithm to efficiently find blobs in the fusion experiments and numerical simulations. We have implemented this algorithm with hybrid MPI/OpenMP and demonstrated the accuracy and efficiency with a set of data from the XGC1 fusion simulations. Our tests

show that we can complete blob detection in a few milliseconds using a cluster at NERSC and achieve linear time scalability. We plan to apply the detection algorithm to experimental measurement data from operating fusion devices. We also plan to develop a blob tracking algorithm based on the proposed method.

Scalable Asynchronous Contact Mechanics with Charm++
Xiang Ni (University of Illinois at Urbana-Champaign)

This poster presents a scalable implementation of the Asynchronous Contact Mechanics (ACM) algorithm, a reliable method to simulate flexible material subject to complex collisions and contact geometries, e.g. cloth simulation in animation. The parallelization of ACM is challenging due to its highly irregular communication pattern, requirement for dynamic load balancing, and extremely fine-grained computations. We make use of Charm++, an adaptive parallel runtime system, to address these challenges and scale ACM to 384 cores for problems with less than 100k vertices. By comparison the previously published shared memory implementation only scales well to about 30 cores. We demonstrate the scalability of our implementation through a number of challenging examples. In particular, for a simulation of a cylindrical rod twisting within a cloth sheet, the simulation time is reduced by 12× from 9 hours on 30 cores to 46 minutes using our implementation on 384 cores of a Cray XC30.

Comparing Decoupled I/O Kernels versus Real Traces in the I/O Analysis of the HACC Scientific Applications on Large-Scale Systems

Sean McDaniel (University of Delaware)

Synthetic benchmarks are often inadequate in comparison to real-world applications as large-scale end-to-end system stressors. Current solutions fail to put systems under significant stress and running real-world applications for testing and tuning is not always feasible because requiring increased resources and time. Lack in adequate solutions for stressing systems has increased efforts in decoupling I/O from applications to be used as benchmarks. Currently there is not a general methodology for extracting I/O and communication codes from scientific applications. In this work we compare the real I/O traces of a cosmology framework the Hardware/Hybrid Accelerated Cosmology Codes (HACC) with the traces of its decoupled I/O kernel (HACC-I/O). We validate results of the decoupled I/O kernel versus the HACC's I/O of simulations to quantify whether the kernel reliably mimics the I/O patterns of the real application. We outline similarities and discrepancies; the latter open exciting research opportunities in the HACC's IO tuning.



Tutorials

The SC Tutorials program is always one of the highlights of the SC Conference, offering attendees a variety of short courses on key topics and technologies relevant to high performance computing, networking, storage, and analysis. Tutorials also provide the opportunity to interact with recognized leaders in the field and to learn about the latest technology trends, theory, and practical techniques. As in years past, tutorial submissions were subjected to a rigorous peer review process. Of the 77 submissions, the Tutorials Committee selected 33 tutorials for presentation.

Tutorials

Sunday, November 16

A Computation-Driven Introduction to Parallel Programming in Chapel

8:30am-12pm

Room: 388

Bradford L. Chamberlain, Greg Titus, Sung-Eun Choi (Cray Inc.)

Chapel (<http://chapel.cray.com>) is an emerging parallel language whose design and development are being led by Cray Inc. in collaboration with members of computing labs, academia, and industry—both domestically and internationally. Chapel aims to vastly improve programmability, generality, and portability compared to current parallel programming models while supporting comparable or improved performance. Chapel's design and implementation are portable and open-source, supporting a wide spectrum of platforms from desktops (Mac, Linux, and Windows) to commodity clusters and large-scale systems developed by Cray and other vendors.

This tutorial will provide an in-depth introduction to Chapel's concepts and features using a computation-driven approach: rather than simply lecturing on individual language features, we will introduce each Chapel concept by studying its use in a real computation taken from a motivating benchmark or proxy application. As time permits, we will also demonstrate Chapel interactively in lieu of a hands-on session. We'll wrap up the tutorial by providing an overview of Chapel status and activities, and by soliciting participants for their feedback to improve Chapel's utility for their parallel computing needs.

A "Hands-On" Introduction to OpenMP

8:30am-5pm

Room: 395

Tim Mattson (Intel Corporation), Mark Bull (Edinburgh Parallel Computing Centre), Mike Pearce (Intel Corporation)

OpenMP is the de facto standard for writing parallel applications for shared memory computers. With multi-core processors in everything from tablets to high-end servers, the need for multithreaded applications is growing, and OpenMP is one of the most straightforward ways to write such programs. In this tutorial, we will cover the core features of the OpenMP 4.0 standard.

This is a hands-on tutorial. We expect you to use your own laptops (with Windows, Linux, or OS/X). You will have access to systems with OpenMP (a remote SMP server), but the best option is for you to load an OpenMP compiler onto your laptops before the tutorial. Information about OpenMP compilers is available at www.openmp.org.

Efficient Parallel Debugging for MPI, Threads and Beyond

8:30am-5pm

Room: 398-99

Matthias S. Müller (Aachen University), David Lecomber (Allinea Software), Tobias Hilbrich (Technische Universität Dresden), Mark O'Connor (Allinea Software), Bronis R. de Supinski (Lawrence Livermore National Laboratory), Ganesh Gopalakrishnan (University of Utah), Joachim Protze (Aachen University)

Parallel software enables modern simulations on high performance computing systems. Defects—or commonly bugs—in these simulations can have dramatic consequences on matters of utmost importance. This is especially true if defects remain unnoticed and silently corrupt results. The correctness of parallel software is a key challenge for simulations in which we can trust. At the same time, the parallelism that enables these simulations also challenges their developers since it gives rise to new sources of defects. We invite attendees to a tutorial that addresses the efficient removal of software defects. The tutorial provides systematic debugging techniques that are tailored to defects that revolve around parallelism. We present leading edge tools that aid developers in pinpointing, understanding and removing defects in MPI, OpenMP, MPI-OpenMP, and further parallel programming paradigms. The tutorial tools include the widely used parallel debugger Allinea DDT, the Intel Inspector XE that reveals data races, and the two MPI correctness tools MUST and ISP. Hands-on examples—make sure to bring a laptop computer—will guide attendees in the use of these tools. We will conclude the tutorial with a discussion and will provide pointers for debugging with paradigms such as CUDA or on architectures such as Xeon Phi.

From "Hello World" to Exascale Using x86, GPUs and Intel Xeon Phi Coprocessors

8:30am-5pm

Room: 396

Robert M. Farber (Blackdog Endeavors, LLC)

Both GPUs and Intel Xeon Phi coprocessors can provide a teraflop/s performance. Working source code will demonstrate how to achieve such high performance using OpenACC, OpenMP, CUDA, and Intel Xeon Phi. Key data structures for GPUs and multi-core, such as low-wait counters, accumulators and massively-parallel stack will be covered. Short understandable examples will walk students from "Hello World" first programs to exascale capable computation via a generic mapping for numerical optimization that demonstrates near-linear scaling on conventional, GPU, and Intel Xeon Phi based leadership-

class supercomputers. Students will work hands-on with code that actually delivers a teraflop/s average performance per device plus MPI code that scales to the largest leadership class supercomputers, and leave with the ability to solve generic optimization problems including data intensive PCA (Principle Components Analysis), NLPCA (Nonlinear Principle Components), plus numerous machine learning and optimization algorithms. A generic framework for data intensive computing will be discussed and provided. Real-time visualization and video processing will also be covered because GPUs make superb big-data visualization platforms.

Hands-On Practical Hybrid Parallel Application Performance Engineering

8:30am-5pm

Room: 383-84-85

Markus Geimer (Forschungszentrum Juelich GmbH), Sameer S. Shende (University of Oregon), Bert Wesarg (Technical University Dresden), Brian J. N. Wylie (Forschungszentrum Juelich GmbH)

This tutorial presents state-of-the-art performance tools for leading-edge HPC systems founded on the Score-P community-developed instrumentation and measurement infrastructure, demonstrating how they can be used for performance engineering of effective scientific applications based on standard MPI, OpenMP, hybrid combination of both, and increasingly common usage of accelerators. Parallel performance tools from the Virtual Institute-High Productivity Supercomputing (VI-HPS) are introduced and featured in hands-on exercises with Scalasca, Vampir, and TAU. We present the complete workflow of performance engineering, including instrumentation, measurement (profiling and tracing, timing and PAPI hardware counters), data storage, analysis, and visualization. Emphasis is placed on how tools are used in combination for identifying performance problems and investigating optimization alternatives. Using their own notebook computers with a provided Linux Live-ISO image containing all of the tools (running within a virtual machine or booted directly from DVD/USB) will help to prepare participants to locate and diagnose performance bottlenecks in their own parallel programs.

How to Analyze the Performance of Parallel Codes 101

8:30am-12pm

Room: 392

Martin Schulz (Lawrence Livermore National Laboratory), James E. Galarowicz (Krell Institute), Mahesh Rajan (Sandia National Laboratories)

Performance analysis is an essential step in the development of HPC codes. It will even gain in importance with the rising complexity of machines and applications that we are seeing today. Many tools exist to help with this analysis, but the user is too often left alone with interpreting the results.

In this tutorial we will provide a practical road map for the performance analysis of HPC codes and will provide users step-by-step advice on how to detect and optimize common performance problems in HPC codes. We will cover both on-node performance and communication optimization and will also touch on threaded and accelerator-based architectures. Throughout this tutorial, we will show live demos using Open|SpeedShop, a comprehensive and easy to use performance analysis tool set, to demonstrate the individual analysis steps. All techniques will, however, apply broadly to any tool and we will point out alternative tools where useful.

Large Scale Visualization with ParaView

8:30am-5pm

Room: 389

Kenneth Moreland, W. Alan Scott (Sandia National Laboratories); David E. DeMarle (Kitware Inc.); Joseph Insley (Argonne National Laboratory), Ollie Lo (Los Alamos National Laboratory), Robert Maynard (Kitware Inc.)

ParaView is a powerful open-source turnkey application for analyzing and visualizing large data sets in parallel. Designed to be configurable, extendible and scalable, ParaView is built upon the Visualization Toolkit (VTK) to allow rapid deployment of visualization components. This tutorial presents the architecture of ParaView and the fundamentals of parallel visualization. Attendees will learn the basics of using ParaView for scientific visualization with hands-on lessons. The tutorial features detailed guidance in visualizing the massive simulations run on today's supercomputers and an introduction to scripting and extending ParaView. Attendees should bring laptops to install ParaView and follow along with the demonstrations.

Linear Algebra Libraries for High-Performance Computing: Scientific Computing with Multicore and Accelerators

8:30am-5pm

Room: 391

Jakub Kurzak (University of Tennessee), Michael Heroux (Sandia National Laboratories), James Demmel (University of California, Berkeley)

Today, desktops with a multicore processor and a GPU accelerator can already provide a TeraFlop/s of performance, while the performance of the high-end systems, based on multicores and accelerators, is already measured in PetaFlop/s. This tremendous computational power can only be fully utilized with the appropriate software infrastructure, both at the low end (desktop, server) and at the high end (supercomputer installation). Most often a major part of the computational effort in scientific and engineering computing goes in solving linear algebra subproblems. After providing a historical overview of legacy software packages, this tutorial surveys the current state of the art numerical libraries for solving problems in lin-

ear algebra, both dense and sparse. MAGMA, (D)PLASMA and Trilinos software packages are discussed in detail. The tutorial also highlights recent advances in algorithms that minimize communication, i.e. data motion, which is much more expensive than arithmetic.

MPI+X: Hybrid Programming on Modern Compute Clusters with Multicore Processors and Accelerators

8:30am-12pm

Room: 393

Rolf Rabenseifner (High Performance Computing Center Stuttgart), Georg Hager (Erlangen Regional Computing Center)

Most HPC systems are clusters of shared memory nodes. Such SMP nodes can be small multi-core CPUs up to large many-core CPUs. Parallel programming may combine the distributed memory parallelization on the node interconnect (e.g., with MPI) with the shared memory parallelization inside of each node (e.g., with OpenMP or MPI-3.0 shared memory). This tutorial analyzes the strengths and weaknesses of several parallel programming models on clusters of SMP nodes. Multi-socket-multi-core systems in highly parallel environments are given special consideration. MPI-3.0 introduced a new shared memory programming interface, which can be combined with inter-node MPI communication. It can be used for direct neighbor accesses similar to OpenMP or for direct halo copies, and enables new hybrid programming models. These models are compared with various hybrid MPI+OpenMP approaches and pure MPI. This tutorial also includes a discussion on OpenMP support for accelerators. Benchmark results are presented for modern platforms such as Intel Xeon Phi and Cray XC30. Numerous case studies and micro-benchmarks demonstrate the performance-related aspects of hybrid programming. The various programming schemes and their technical and performance implications are compared. Tools for hybrid programming such as thread/process placement support and performance analysis are presented in a “how-to” section. Details: <https://fs.hlr.de/projects/rabenseifner/publ/SC2014-hybrid.html>

Parallel Computing 101

8:30am-5pm

Room: 386-87

Quentin F. Stout, Christiane Jablonowski (University of Michigan)

This tutorial provides a comprehensive overview of parallel computing, emphasizing those aspects most relevant to the user. It is suitable for new users, managers, students and anyone seeking an overview of parallel computing. It discusses software and hardware/software interaction, with an emphasis on standards, portability, and systems that are widely available.

The tutorial surveys basic parallel computing concepts, using examples selected from multiple engineering and scientific problems. These examples illustrate using MPI on distributed memory systems, OpenMP on shared memory systems, MPI+OpenMP on hybrid systems, GPU programming, and Hadoop on big data. It discusses numerous parallelization and load balancing approaches, and software engineering and performance improvement aspects, including the use of state-of-the-art tools.

The tutorial helps attendees make intelligent decisions by covering the primary options that are available, explaining how they are used and what they are most suitable for. Extensive pointers to the literature and web-based resources are provided to facilitate follow-up studies.

Parallel Programming in Modern Fortran

8:30am-5pm

Room: 397

Karla Morris (Sandia National Laboratories), Damian Rouson (Sourcery, Inc.), Salvatore Filippone (University of Rome Tor Vergata)

User surveys from HPC centers in the U.S. and Europe consistently show Fortran predominating, but most users describe their programming skills as self-taught and most continue to use older versions of the language. The increasing compiler support for the parallel programming features of Fortran 2008 makes the time ripe to offer instruction on these features to the HPC community. This tutorial teaches single-program, multiple-data (SPMD) programming with Fortran 2008 coarrays. We also introduce Fortran’s loop concurrency and pure procedure features and demonstrate their use in asynchronous expression evaluation for partial differential equation (PDE) solvers. We incorporate other language features, including object-oriented (OO) programming, when they support our chief aim of teaching parallel programming. In particular, we demonstrate OO design patterns that enable hybrid CPU/GPU calculations on sparse matrices in the Parallel Sparse Basic Linear Algebra Subroutines (PSBLAS) library. The students will obtain hands-on experience with parallel programming in modern Fortran through the use of virtual machines.

Programming the Xeon Phi

8:30am-5pm

Room: 394

*Jerome Vienne, Kent Milfeld, Lars Koesterke, Dan Stanzione
(Texas Advanced Computing Center, University of Texas at Austin)*

The use of heterogeneous architectures in HPC at the large scale has become increasingly common over the past few years. One new technology for HPC is the Intel Xeon Phi co-processor, also known as the MIC. The Xeon Phi is x86 based, hosts its own Linux OS, and is capable of running most codes with little porting effort. However, the MIC architecture has significant features that are different from that of current x86 CPUs. Attaining optimal performance requires an understanding of possible execution models and the architecture.

This tutorial is designed to introduce attendees to the MIC architecture in a practical manner. Experienced C/C++ and Fortran programmers will be introduced to techniques essential for utilizing the MIC architecture efficiently. Multiple lectures and hands-on exercises will be used to acquaint attendees with the MIC platform and to explore the different execution modes as well as parallelization and optimization through example testing and reports. All exercises will be executed on the Stampede system at the Texas Advanced Computing Center (TACC). Stampede features more than 2PF of performance using 100,000 Intel Xeon E5 cores and an additional 7+ PF of performance from more than 6,400 Xeon Phi.

SciDB - Manage and Analyze Terabytes of Array Data

8:30am-5pm

Room: 390

*Yushu Yao (Lawrence Berkeley National Laboratory),
Alex Poliakov (Paradigm4), Jeremy Kepner (Massachusetts
Institute of Technology Lincoln Laboratory), Lisa Gerhardt (Lawrence Berkeley National Laboratory), Marilyn Matz (Paradigm4)*

With the emergence of the Internet of Everything in the commercial and industrial worlds and with the advances in device and instrument technologies in the science world, there is an urgent need for data scientists and scientists to be able to work more easily with extremely large and diverse data sets. SciDB is an open-source analytical database for scalable complex analytics on very large array or multi-structured data from a variety of sources, programmable from R and Python. It runs on HPC, commodity hardware grids, or in a cloud. We present an overview of SciDB's array data model, programming and query interfaces, and math library. We will demonstrate how to set up a SciDB cluster, ingest data from various standard file formats, design schema, and analyze data on two use cases: one from computational genomics and one using satellite imagery of the earth. This tutorial will help computational scientists learn how to do interactive exploratory data mining and

analytics on terabytes of data. During the tutorial, attendees will have an opportunity to describe their own data and key operations needed for analysis. The presenters can guide them through implementing their use case in SciDB.

PGAS and Hybrid MPI+PGAS Programming Models on Modern HPC Clusters

1:30pm-5pm

Room: 393

Dhabaleswar K. Panda, Khaled Hamidouche (Ohio State University)

Multi-core processors, accelerators (GPUs/MIC) and high-performance interconnects with RDMA are shaping the architecture for next generation exascale clusters. Efficient programming models to design applications on these systems are still evolving. Partitioned Global Address Space (PGAS) models provide an attractive alternative to the traditional MPI model, owing to their easy to use shared memory abstractions and light-weight one-sided communication. Hybrid MPI+PGAS models are gaining attention as a possible solution to programming exascale systems. They help MPI applications to take advantage of PGAS models, without paying the prohibitive cost of re-designing complete applications. They also enable hierarchical design of applications using different models to suite modern architectures. In this tutorial, we provide an overview of the research and development taking place and discuss associated opportunities and challenges as we head toward exascale. We start with an in-depth overview of modern system architectures with multi-core processors, accelerators and high-performance interconnects. We present an overview of UPC and OpenSHMEM. We introduce MPI+PGAS hybrid programming models and highlight their advantages and challenges. We examine the challenges in designing high-performance UPC, OpenSHMEM and unified MPI+UPC/OpenSHMEM runtimes. We present application case-studies to demonstrate the productivity and performance of MPI+PGAS models, using the publicly available MVAPICH2-X software package.

Practical Fault Tolerance on Today's Supercomputing Systems

1:30pm-5pm

Room: 388

*Kathryn Mohror (Lawrence Livermore National Laboratory),
Nathan DeBardeleben (Los Alamos National Laboratory),
Eric Roman (Lawrence Berkeley National Laboratory),
Laxmikant Kale (University of Illinois at Urbana-Champaign)*

The failure rates on high performance computing systems are increasing with increasing component count. Applications running on these systems currently experience failures on the order of days; however, on future systems, predictions of failure rates range from minutes to hours. Developers need

to defend their application runs from losing valuable data by using fault tolerant techniques. These techniques range from changing algorithms, to checkpoint and restart, to programming model-based approaches. In this tutorial, we will present introductory material for developers who wish to learn fault tolerant techniques available on today's systems. We will give background information on the kinds of faults occurring on today's systems and trends we expect going forward. Following this, we will give detailed information on several fault tolerant approaches and how to incorporate them into applications. Our focus will be on scalable checkpoint and restart mechanisms and programming model-based approaches.

Scaling I/O Beyond 100,000 Cores Using ADIOS

1:30pm-5pm

Room: 392

Norbert Podhorszki, Scott Klasky, Qing Liu (Oak Ridge National Laboratory)

As concurrency and complexities continue to increase on high-end machines, in terms of both the number of cores and the level of storage hierarchy, managing I/O efficiently becomes more and more challenging. As we are moving forward, one of the major roadblocks to exascale is how to manipulate, write, read big datasets quickly and efficiently on high-end machines. In this tutorial we will demonstrate I/O practices and techniques that are crucial to achieve high performance on 100,000+ cores. Part I of this tutorial will introduce parallel I/O and the ADIOS framework to the audience. Specifically we will discuss the concept of ADIOS I/O abstraction, the binary-packed file format, and I/O methods along with the benefits to applications. Since 1.4.1, ADIOS can operate on both files and data streams. Part II will include a session on how to write/read data and how to use different I/O componentizations inside of ADIOS. Part III will show users how to take advantage of the ADIOS framework to do topology-aware data movement, compression and data staging/streaming using DIMES/FLEXPATh.

Monday, November 17

Advanced MPI Programming

8:30am-5pm

Room: 392

Pavan Balaji (Argonne National Laboratory), William Gropp (University of Illinois at Urbana-Champaign), Torsten Hoeftler (ETH Zurich), Rajeev Thakur (Argonne National Laboratory)

The vast majority of production-parallel scientific applications today use MPI and run successfully on the largest systems in the world. For example, several MPI applications are running at full scale on the Sequoia system (on 1.6 million cores) and achieving 12 to 14 petaflops/s of sustained performance. At the same time, the MPI standard itself is evolving (MPI-3 was released in late 2012) to address the needs and challenges of future extreme-scale platforms as well as applications. This tutorial will cover several advanced features of MPI, including new MPI-3 features, that can help users program modern systems effectively. Using code examples based on scenarios found in real applications, we will cover several topics including efficient ways of doing 2D and 3D stencil computation, derived datatypes, one-sided communication, hybrid (MPI + shared memory) programming, topologies and topology mapping, and neighborhood and nonblocking collectives. Attendees will leave the tutorial with an understanding of how to use these advanced features of MPI and guidelines on how they might perform on different platforms and architectures.

Advanced OpenMP: Performance and 4.0 Features

8:30am-5pm

Room: 397

Christian Terboven (RWTH Aachen University), Ruud van der Pas (Oracle), Eric J. Stotzer (Texas Instruments), Bronis R. de Supinski (Lawrence Livermore National Laboratory), Michael Klemm (Intel Corporation)

With the increasing prevalence of multicore processors, shared-memory programming models are essential. OpenMP is a popular, portable, widely supported and easy-to-use shared-memory model. Developers usually find OpenMP easy to learn. However, they are often disappointed with the performance and scalability of the resulting code. This disappointment stems not from shortcomings of OpenMP but rather with the lack of depth with which it is employed. Our "Advanced OpenMP Programming" tutorial addresses this critical need by exploring the implications of possible OpenMP parallelization strategies, both in terms of correctness and performance.

While we quickly review the basics of OpenMP programming, we assume attendees understand basic parallelization concepts and will easily grasp those basics. We focus on performance aspects, such as data and thread locality on NUMA

architectures, false sharing, and exploitation of vector units. We discuss language features in-depth, with emphasis on advanced features like tasking and those recently added to OpenMP 4.0 such as cancellation. We close with the presentation of the new directives for attached compute accelerators.

Debugging and Performance Tools for MPI and OpenMP 4.0 Applications for CPU and Accelerators/Coprocessors

8:30am-5pm

Room: 398-99

Sandra Wienke (RWTH Aachen University), Mike Ashworth (STFC Daresbury Laboratory), Damian Alvarez (Forschungszentrum Juelich), Vince Betto (University of Tennessee, Knoxville), Chris Gottbrath, Nikolay Piskun (Rogue Wave Software)

With HPC trends heading towards increasingly heterogeneous solutions, scientific developers face challenges adapting software to leverage these new architectures. For instance, many systems feature nodes that couple multi-core processors with GPU-based computational accelerators, like the NVIDIA Kepler, or many-core coprocessors, like the Intel Xeon Phi. In order to effectively utilize these systems, programmers need to leverage an extreme level of parallelism in applications. Developers also need to juggle multiple programming paradigms including MPI, OpenMP, CUDA, and OpenACC.

This tutorial provides in-depth exploration of parallel debugging and optimization focused on techniques that can be used with accelerators and coprocessors. We cover debugging techniques such as grouping, advanced breakpoints and barriers, and MPI message queue graphing. We discuss optimization techniques like profiling, tracing, and cache memory optimization with tools such as Vampir, Scalasca, Tau, CrayPAT, Vtune and the NVIDIA Visual Profiler. Participants have the opportunity to do hands-on GPU and Intel Xeon Phi debugging and profiling. Additionally, the OpenMP 4.0 standard will be covered which introduces novel capabilities for both Xeon Phi and GPU programming. We will discuss peculiarities of that specification with respect to error finding and optimization. A laptop will be required for hands-on sessions.

Fault-tolerance for HPC: Theory and Practice

8:30am-5pm

Room: 388

Thomas Hérault (University of Tennessee, Knoxville), Yves Robert (ENS Lyon); George Bosilca, Aurélien Bouteiller (University of Tennessee, Knoxville)

Resilience is a critical issue for large-scale platforms. This tutorial provides a comprehensive survey of fault-tolerant techniques for high-performance computing, with a fair balance between practice and theory. It is organized along four main topics: (1) An overview of failure types (software/hardware,

transient/fail-stop) and typical probability distributions (Exponential, Weibull, Log-Normal); (2) General-purpose techniques, which include several checkpoint and rollback recovery protocols, replication, prediction and silent error detection; (3) Application-specific techniques, such as ABFT for grid-based algorithms or fixed-point convergence for iterative applications; and (4) Practical deployment of fault tolerant techniques with User Level Fault Mitigation (a proposed MPI extension to the MPI forum). Relevant examples based on widespread computational solver routines will be protected with a mix of checkpoint-restart and advanced recovery techniques in a hands-on session.

The tutorial is open to all SC14 attendees who are interested in the current status and expected promise of fault-tolerant approaches for scientific applications. There are no audience prerequisites: background will be provided for all protocols and probabilistic models. However, basic knowledge of MPI will be helpful for the hands-on session.

In Situ Data Analysis and Visualization with ParaView Catalyst

8:30am-12pm

Room: 394

David H. Rogers (Los Alamos National Laboratory), Jeffrey Mauldin (Sandia National Laboratories), Thierry Carrard (French Alternative Energies and Atomic Energy Commission), Andrew Bauer (Kitware, Inc.)

As supercomputing moves towards exascale, scientists, engineers and medical researchers will look for efficient and cost effective ways to enable data analysis and visualization for the products of their computational efforts. The ‘exa’ metric prefix stands for quintillion, and the proposed exascale computers would approximately perform as many operations per second as 50 million laptops. Clearly, typical spatial and temporal data reduction techniques employed for post processing will not yield desirable results where reductions of 10e3, 10e6, or 10e9 may still produce petabytes, terabytes or gigabytes of data to transfer or store. Since transferring or storing data may no longer be viable for many simulation applications, data analysis and visualization must now be performed in situ. ParaView Catalyst is an open-source data analysis and visualization library, which aims to reduce IO by tightly coupling simulation, data analysis and visualization codes. This tutorial presents the architecture of ParaView Catalyst and the fundamentals of in situ data analysis and visualization. Attendees will learn the basics of using ParaView Catalyst with hands-on exercises. The tutorial features detailed guidance in implementing C++, Fortran and Python examples. Attendees should bring laptops to install a VirtualBox image and follow along with the demonstrations.

InfiniBand and High-Speed Ethernet for Dummies**8:30am-12pm****Room: 389***Dhabaleswar K. Panda, Hari Subramoni (Ohio State University)*

InfiniBand (IB) and High-speed Ethernet (HSE) technologies are generating a lot of excitement towards building next generation High-End Computing (HEC) systems including clusters, datacenters, file systems, storage, cloud computing and Big Data (Hadoop, HBase and Memcached) environments. RDMA over Converged Enhanced Ethernet (RoCE) technology is also emerging. This tutorial will provide an overview of these emerging technologies, their offered architectural features, their current market standing, and their suitability for designing HEC systems. It will start with a brief overview of IB and HSE. In-depth overview of the architectural features of IB and HSE (including iWARP and RoCE), their similarities and differences, and the associated protocols will be presented. Next, an overview of the emerging OpenFabrics stack which encapsulates IB, HSE and RoCE in a unified manner will be presented. Hardware/software solutions and the market trends behind IB, HSE and RoCE will be highlighted. Finally, sample performance numbers of these technologies and protocols for different environments will be presented.

Introducing R: From Your Laptop to HPC and Big Data**8:30am-12pm****Room: 391***George Ostrouchov (Oak Ridge National Laboratory),
Drew Schmidt (University of Tennessee)*

The R language has been called the “lingua franca” of data analysis and statistical computing and is quickly becoming the de facto standard for analytics. This tutorial will introduce attendees to the basics of the R language with a focus on its recent high performance extensions enabled by the “Programming with Big Data in R” (pbdR) project. Although R had a reputation for lacking scalability, our experiments with pbdR have easily scaled to 50 thousand cores. No background in R is assumed but even R veterans will benefit greatly from the session. We will cover only those basics of R that are needed for the HPC portion of the tutorial. The tutorial is very much example-oriented, with many opportunities for the engaged attendee to follow along. Examples on real data will utilize common data analytics techniques, such as principal components analysis and cluster analysis.

Node-Level Performance Engineering**8:30am-5pm****Room: 390***Georg Hager, Jan Treibig, Gerhard Wellein (Erlangen Regional Computing Center)*

The advent of multi- and manycore chips has led to a further opening of the gap between peak and application performance for many scientific codes. This trend is accelerating as we move from petascale to exascale. Paradoxically, bad node-level performance helps to “efficiently” scale to massive parallelism, but at the price of increased overall time to solution. If the user cares about time to solution on any scale, optimal performance on the node level is often the key factor. We convey the architectural features of current processor chips, multiprocessor nodes, and accelerators, as far as they are relevant for the practitioner. Peculiarities like SIMD vectorization, shared vs. separate caches, bandwidth bottlenecks, and ccNUMA characteristics are introduced, and the influence of system topology and affinity on the performance of typical parallel programming constructs is demonstrated. Performance engineering and performance patterns are suggested as powerful tools that help the user understand the bottlenecks at hand and to assess the impact of possible code optimizations. A cornerstone of these concepts is the roofline model, which is described in detail, including useful case studies, limits of its applicability, and possible refinements.

OpenACC: Productive, Portable Performance on Hybrid Systems using High-Level Compilers and Tools**8:30am-5pm****Room: 396***Luiz DeRose, Alistair Hart, Heidi Poxon, James Beyer (Cray Inc.)*

Portability and programming difficulty are two critical hurdles in generating widespread adoption of accelerated computing in high performance computing. The dominant programming frameworks for accelerator-based systems (CUDA and OpenCL) offer the power to extract performance from accelerators, but with extreme costs in usability, maintenance, development, and portability. To be an effective HPC platform, hybrid systems need a high-level programming environment to enable widespread porting and development of applications that run efficiently on either accelerators or CPUs. In this hands-on tutorial we present the high-level OpenACC parallel programming model for accelerator-based systems, demonstrating compilers, libraries, and tools that support this cross-vendor initiative. Using personal experience in porting large-scale HPC applications, we provide development guidance, practical tricks, and tips to enable effective and efficient use of these hybrid systems, both in terms of runtime and energy efficiency.

OpenCL: A Hands-On Introduction

8:30am-5pm

Room: 395

Tim Mattson (Intel Corporation), Alice Koniges (Lawrence Berkeley National Laboratory), Simon McIntosh-Smith (University of Bristol)

OpenCL is an open standard for programming heterogeneous parallel computers composed of CPUs, GPUs and other processors. OpenCL consists of a framework to manipulate the host CPU and one or more compute devices plus a C-based programming language for writing programs for the compute devices. Using OpenCL, a programmer can write parallel programs that harness all of the resources of a heterogeneous computer.

In this hands-on tutorial, we introduce OpenCL using the more accessible C++ API. The tutorial format will be a 50/50 split between lectures and exercises. Students will use their own laptops (Windows, Linux or OS/X) and log into a remote server running an OpenCL platform. Alternatively, students can load OpenCL onto their own laptops prior to the course (Intel, AMD and NVIDIA provide OpenCL SDKs. Apple laptops with X-code include OpenCL by default. Be sure to configure X-code to use the command line interface).

The last segment of the tutorial will be spent visiting the “OpenCL zoo”; a diverse collection of OpenCL conformant devices. Tutorial attendees will run their own programs on devices in the zoo to explore performance portability of OpenCL. The zoo should include a mix of CPU, GPU, FPGA, mobile, and DSP devices.

Parallel I/O in Practice

8:30am-5pm

Room: 386-87

Robert J. Latham, Robert Ross (Argonne National Laboratory); Brent Welch (Google), Katie Antypas (National Energy Research Scientific Computing Center)

I/O on HPC systems is a black art. This tutorial sheds light on the state-of-the-art in parallel I/O and provides the knowledge necessary for attendees to best leverage I/O resources available to them. We cover the entire I/O software stack from parallel file systems at the lowest layer, to intermediate layers (such as MPI-IO), and finally high-level I/O libraries (such as HDF-5). We emphasize ways to use these interfaces that result in high performance. Benchmarks on real systems are used throughout to show real-world results.

This tutorial first discusses parallel file systems in detail (PFSs). We cover general concepts and examine three examples: GPFS, Lustre, and PanFS. We examine the upper layers of the I/O stack, covering POSIX I/O, MPI-IO, Parallel netCDF, and

HDF5. We discuss interface features, show code examples, and describe how application calls translate into PFS operations. Finally we discuss I/O best practice.

Parallel Programming with Charm++

8:30am-12pm

Room: 393

Laxmikant Kale, Michael Robson (University of Illinois at Urbana-Champaign), Nikhil Jain (University of Illinois at Urbana-Champaign)

There are several challenges in programming applications for large supercomputers: exposing concurrency, data movement, load imbalance, heterogeneity, variations in application’s behavior, system failures, etc. Addressing these challenges requires more emphasis on the following important concepts during application development: overdecomposition, asynchrony, migratability, and adaptivity. At the same time, the runtime systems (RTS) will need to become introspective and provide automated support for several tasks, e.g. load balancing, that currently burden the programmer.

This tutorial is aimed at exposing the attendees to the above mentioned concepts. We will present details on how a concrete implementation of these concepts, in synergy with an introspective RTS, can lead to development of applications that scale irrespective of the rough landscape. We will focus on Charm++ as the programming paradigm that encapsulates these ideas, and use examples from real world applications to further the understanding.

Charm++ provides an asynchronous, message-driven programming model via migratable objects and an adaptive RTS that guides execution. It automatically overlaps communication, balances loads, tolerates failures, checkpoints for split-execution, interoperates with MPI, and promotes modularity while allowing programming in C++. Several widely used Charm++ applications thrive in computational science domains including biomolecular modeling, cosmology, quantum chemistry, epidemiology, and stochastic optimization.

Python in HPC

8:30am-5pm

Room: 383-84-85

Andy Terrel (Continuum Analytics), Matt Knepley (University of Chicago), Kurt Smith (Enthought, Inc), Matthew Turk (Columbia University)

The Python ecosystem empowers the HPC community with a stack of tools that are not only powerful but a joy to work with. It is consistently one of the top languages in HPC with a growing vibrant community of open source tools. Proven to scale on the world’s largest clusters, it is a language that has continued to innovate with a wealth of new data tools.

This tutorial will survey the state of the art tools and techniques used by HPC Python experts throughout the world. The first half of the day will include an introduction to the standard toolset used in HPC Python and techniques for speeding Python and using legacy codes by wrapping Fortran and C. The second half of the day will include discussion on using Python in a distributed workflow via MPI and tools for handling large scale visualizations.

Students should be familiar with basic Python syntax, we recommend the Python 2.7 tutorial on python.org. We will include hands-on demonstrations of building simulations, wrapping low-level code, executing on a cluster via MPI, and use of visualization tools. Examples for a range of experience levels will be provided.

Designing and Using High-End Computing Systems with InfiniBand and High-Speed Ethernet

1:30pm-5pm

Room: 389

Dhabaleswar K. Panda, Hari Subramoni (Ohio State University)

As InfiniBand (IB) and High-Speed Ethernet (HSE) technologies mature, they are being used to design and deploy different kinds of High-End Computing (HEC) systems: HPC clusters with accelerators (GPUs and MIC) supporting MPI and PGAS (UPC and OpenSHMEM), Storage and Parallel File Systems, Cloud Computing with Virtualization, Big Data systems with Hadoop (HDFS, MapReduce and HBase), Multi-tier Datacenters with Web 2.0 (memcached) and Grid Computing systems. These systems are bringing new challenges in terms of performance, scalability, and portability. Many scientists, engineers, researchers, managers and system administrators are becoming interested in learning about these challenges, approaches being used to solve these challenges, and the associated impact on performance and scalability. This tutorial will start with an overview of these systems and a common set of challenges being faced while designing these systems. Advanced hardware and software features of IB and HSE and their capabilities to address these challenges will be emphasized. Next, case studies focusing on domain-specific challenges in designing these systems (including the associated software stacks), their solutions and sample performance numbers will be presented. The tutorial will conclude with a set of demos focusing on RDMA programming, network management infrastructure and tools to effectively use these systems.

Effective HPC Visualization and Data Analysis using VisIt

1:30pm-5pm

Room: 391

Cyrus Harrison (Lawrence Livermore National Laboratory), Jean M. Favre (Swiss National Supercomputing Center), Brad Whitlock (Intelligent Light), David Pugmire (Oak Ridge National Laboratory), Rob Sisneros (National Center for Supercomputing Applications, University of Illinois at Urbana-Champaign)

Visualization and data analysis are an essential component of the scientific discovery process. Scientists and businesses running HPC simulations leverage visualization and analysis tools for data exploration, quantitative analysis, visual debugging, and communication of results. This half-day tutorial will provide attendees with a practical introduction to mesh-based HPC visualization and analysis using VisIt, an open source parallel scientific visualization and data analysis platform. We will provide a foundation in basic HPC visualization practices and couple this with hands-on experience creating visualizations and analyzing data.

This tutorial includes: 1) an introduction to visualization techniques for mesh-based simulations ; 2) a guided tour of VisIt; and 3) hands-on demonstrations of end-to-end visualizations of both a water flow simulation and blood flow (aneurysm) simulation.

This tutorial builds on the past success of VisIt tutorials, updated and anchored with new concrete use cases. Attendees will gain practical knowledge and recipes to help them effectively use VisIt to analyze data from their own simulations.

Enhanced Campus Bridging via a Campus Data Service Using Globus and the Science DMZ

1:30pm-5pm

Room: 394

Steve Tuecke, Vas Vasiliadis (University of Chicago), Raj Ketimuthu (Argonne National Laboratory)

Existing campus data services are limited in their reach and utility due, in part, to unreliable tools and a wide variety of storage systems with sub-optimal user interfaces. An increasingly common solution to campus bridging comprises Globus operating within the Science DMZ, enabling reliable, secure file transfer and sharing, while optimizing use of existing high-speed network connections and campus identity infrastructures. Attendees will be introduced to Globus and have the opportunity for hands-on interaction installing and configuring the basic components of a campus data service. We will also describe how newly developed Globus services for public cloud storage integration and metadata management may be used as the basis for a campus publication system that meets

an increasingly common need at many campus libraries. The tutorial will help participants answer these questions: What services can I offer to researchers for managing large datasets more efficiently? How can I integrate these services into existing campus computing infrastructure? What role can the public cloud play (and how does a service like Globus facilitate its integration)? How should such services be delivered to minimize the impact on my infrastructure? What issues should I expect to face (e.g. security), and how should I address them?

Introductory and Advanced OpenSHMEM Programming

1:30pm-5pm

Room: 393

Tony Curtis, Swaroop Pophale (University of Houston), Oscar Hernandez, Pavel Shamis (Oak Ridge National Laboratory); Deepak Eachempati, Aaron Welch (University of Houston)

As high performance computing systems become larger better mechanisms are required to communicate and coordinate processes within or across nodes. Communication libraries like OpenSHMEM can leverage hardware capabilities such as remote directory memory access (RDMA) to hide latencies through one-sided data transfers. OpenSHMEM API provides data transfer routines, collective operations such as reductions and concatenations, synchronizations, and memory management.

In this tutorial, we will present an introductory course on OpenSHMEM, its current state and the community's future plans. We will show how to use OpenSHMEM to add parallelism to programs via an exploration of its core features. We will demonstrate simple techniques to port applications to run at scale while improving the program performance using OpenSHMEM, and discuss how to migrate existing applications that use message passing techniques to equivalent OpenSHMEM programs that run more efficiently. We will then present a new low-level communication library called Unified Common Communication Substrate (UCCS) and the second part of the tutorial will focus on how to use the OpenSHMEM-UCCS implementation on applications. UCCS is designed to sit underneath any message passing or PGAS user-oriented language or library.





SC14
New Orleans, LA | **hpc**
matters.

Visualization & Data Analytics

Visualization & Data Analytics

This year we present a new format for the Visualization and Data Analytics Showcase Program, which provides a forum for the year's most instrumental movies in HPC. Six finalists will compete for the Best Visualization Award, and each finalist will present his or her movie during a dedicated session at SC14 in a 15-minute presentation. Movies are judged based on how their movie illuminates science, by the quality of the movie, and for innovations in the process used for creating the movie.

Visualization & Data Analytics Showcase

Wednesday, November 19

Visualization & Data Analytics Showcase

Chair: Hank Childs (University of Oregon and Lawrence Berkeley National Laboratory)

10:30am-12pm

Room: 386-87

Visualization of Energy Conversion Processes in a Light Harvesting Organelle at Atomic Detail

Melih Sener, John E. Stone, Angela Barragan, Abhishek Singharoy, Ivan Teo, Kirby L. Vandivort, Barry Isralewitz, Bo Liu, Boon Chong Goh, James C. Phillips (University of Illinois at Urbana-Champaign), Lena F. Kourkoutis (Cornell University), C. Neil Hunter (University of Sheffield), Klaus Schulten (University of Illinois at Urbana-Champaign)

The cellular process responsible for providing energy for most life on Earth, namely, photosynthetic light-harvesting, requires the cooperation of hundreds of proteins across an organelle, involving length and time scales spanning several orders of magnitude over quantum and classical regimes. Simulation and visualization of this fundamental energy conversion process pose many unique methodological and computational challenges. We present, in an accompanying movie, light-harvesting in the photosynthetic apparatus found in purple bacteria, the so-called chromatophore. The movie is the culmination of three decades of modeling efforts, featuring the collaboration of theoretical, experimental, and computational scientists. We describe the techniques that were used to build, simulate, analyze, and visualize the structures shown in the movie, and we highlight cases where scientific needs spurred the development of new parallel algorithms that efficiently harness GPU accelerators and petascale computers.

In Situ MPAS-Ocean Image-Based Visualization

James Aherns, Sebastien Jourdain, Patrick O'Leary (Kitware, Inc.), John Patchett, David Rogers, Patricia Fasel (Los Alamos National Laboratory), Andrew Bauer (Kitware, Inc.), Mark Petersen (Los Alamos National Laboratory), Francesca Samsel (University of Texas at Austin), Benjamin Boeckel (Kitware, Inc.)

Due to power and I/O constraints associated with extreme scale scientific simulations, in situ visualization and analysis will become a critical component to scientific discovery. The options for extreme scale visualization and analysis are often presented as a stark contrast: write files to disk for interactive, exploratory analysis, or perform in situ analysis to save data products about phenomena that a scientist knows about in advance. In this video demonstrating large-scale visualization

of MPAS-Ocean simulations, we leveraged a third option based on ParaView Cinema, which is a novel framework for highly interactive, image-based in situ visualization and analysis that promotes exploration.

Visualizations of Molecular Dynamics Simulations of High-Performance Polycrystalline Structural Ceramics

Christopher Lewis, Miguel Valenciano (High Performance Computing Modernization Program Data Analysis and Assessment Center), Charles Cornwell (U.S. Army Engineer Research and Development Center)

The Data Analysis and Assessment Center (DAAC) serves the needs of the Department of Defense (DOD) High Performance Computing Modernization Program (HPCMP) scientists by facilitating the analysis of an ever increasing volume and complexity of data. Dr. Charles Cornwell is a research scientist and HPCMP user who ran nanoscale molecular dynamics simulations using Large-scale Atomic/Molecular Massively Parallel Simulator code (LAMMPS) from Sandia National Laboratories. The largest of the resulting data contained over 15 million atoms. The DAAC developed a novel method to visualize this unusually large-scale and complex data. All of the molecular dynamics simulations and the analytics were processed using High Performance Computing (HPC) resources available to users of the DOD HPCMP.

Investigating Flow-Structure Interactions in Cerebral Aneurysms

Joseph A. Insley (Argonne National Laboratory), Paris Perdikaris (Brown University), Leopold Grinberg (IBM Corporation), Yue Yu (Brown University), Michael E. Papka (Argonne National Laboratory), George Em. Karniadakis (Brown University)

Modeling flow-structure interaction (FSI) in biological systems, where elastic tissues and fluids have almost identical density (low mass ratio problem), is one of the most challenging problems in computational mechanics. One of the major challenges in FSI simulations at low mass ratio is the stability of semi-implicit coupled fluid-structure solvers. Based on our recent advances in developing stable numerical algorithms we have developed a solver capable of tackling problems with mass ratios approaching zero, enabling high-resolution FSI simulations of biological systems, such as blood flow in compliant arteries. Our coupled solver, NekTar, is based on high-order spectral/hp element discretization and iterative coupling between the fluid and structure (solid) domains. Realistic FSI simulations produce very large and complex data sets, yielding the need for parallel data processing and visualization. Here we present our recent advancements in developing an interactive visualization tool which now enables the visualization of such FSI

simulation data. Specifically, we present a ParaView-NekTar interface that couples the ParaView visualization engine with NekTar's parallel libraries, which are employed for the calculation of derived fields in both the fluid and solid domains with spectral accuracy. This interface significantly facilitates the visualization of complex structures under large deformations. The animation of the fluid and structure data is synchronized in time, while the ParaView-NekTar interface enables the visualization of different fields to be superimposed, e.g. fluid jet and structural stress, to better understand the interactions in this multiphysics environment.

Visualization of a Simulated Long-Track EF5 Tornado Embedded Within a Supercell Thunderstorm

Leigh Orf (Central Michigan University), Robert Wilhelmson (University of Illinois at Urbana-Champaign), Louis Wicker (National Severe Storms Laboratory)

Tornadoes are one of nature's most destructive forces, creating winds that can exceed 300 miles per hour. The sheer destructive power of the strongest class of tornado (EF5) makes these tornadoes the subject of active research. However, very little is currently known about why some supercells produce long-track (a long damage path) EF5 tornadoes, while other storms in similar environments produce short-lived, weak tornadoes, or produce no tornado at all.

In this work we visualize cloud model simulation data of a supercell thunderstorm that produces a long-track EF5 tornado. Several obstacles needed to be overcome in order to produce the visualization of this simulation, including managing hundreds of TB of model I/O, interfacing the model output format to a high-quality visualization tool, choosing effective visualization parameters, and, most importantly, actually creating a simulation where a long-track EF5 tornado occurs within the model, which only recently has been accomplished.

A Global Perspective of Atmospheric CO₂ Concentrations

William Putman, Lesley, Anton Darmenov, Arlindo daSilva (National Aeronautics and Space Administration Goddard Space Flight Center)

Carbon dioxide (CO₂) is the most important greenhouse gas affected by human activity. About half of the CO₂ emitted from fossil fuel combustion remains in the atmosphere, contributing to rising temperatures, while the other half is absorbed by natural land and ocean carbon reservoirs. Despite the importance of CO₂, many questions remain regarding the processes that control these fluxes and how they may change in response to a changing climate. The Orbiting Carbon Observatory-2 (OCO-2), launched on July 2, 2014, is NASA's first satellite mission designed to provide the global view of atmospheric CO₂ needed to better understand both human emissions and natural fluxes.

This visualization shows how column CO₂ mixing ratio, the quantity observed by OCO-2, varies throughout the year. By observing spatial and temporal gradients in CO₂ like those shown, OCO-2 data will improve our understanding of carbon flux estimates. But, CO₂ observations can't do that alone. This visualization also shows that column CO₂ mixing ratios are strongly affected by large-scale weather systems.

A high-resolution (7-km) non-hydrostatic global mesoscale simulation using the Goddard Earth Observing System (GEOS-5) model produces the CO₂ concentrations and weather systems in this visualization. This 7-km GEOS-5 Nature Run product will provide synthetic observations for missions like OCO-2. In order to fully understand carbon flux processes, OCO-2 observations and atmospheric models will work closely together to determine when and where observed CO₂ came from. Together, the combination of high-resolution observations and model simulations will guide climate models towards more reliable predictions of future conditions.



Workshops

SC14 includes 35 full-day and half-day workshops that complement the overall Technical Program events, expand the knowledge base of its subject area, and extend its impact by providing greater depth of focus. These workshops are geared toward providing interaction and in-depth discussion of stimulating topics of interest to the HPC community.

New programs highlight innovative technologies in HPC, alongside the traditional program elements that our community relies upon to stay abreast of the complex, changing landscape of HPC. Together these are the elements that make SC the most informative, exciting, and stimulating technical program in HPC!

Workshops

Sunday, November 16

5th SC Workshop on Big Data Analytics: Challenges and Opportunities

9am-5:30pm

Room: 286-87

Ranga Raju Vatsavai (Oak Ridge National Laboratory), Scott Klasky (Oak Ridge National Laboratory), Manish Parashar (Rutgers University)

The recent decade has witnessed a data explosion, and petabyte-sized data archives are not uncommon any more. It is estimated that organizations with high end computing (HEC) infrastructures and data centers are doubling the amount of data that they are archiving every year. On the other hand, computing infrastructures are becoming more heterogeneous. The first four workshops held during SC10-SC13 were a great success. Continuing on this success, we propose to broaden the topic of this workshop with an emphasis on novel middle-ware (e.g., in situ) infrastructures that facilitate efficient data analytics on big data. The proposed workshop intends to bring together researchers, developers, and practitioners from academia, government, and industry to discuss new and emerging trends in high end computing platforms, programming models, middleware and software services, and outline the data mining and knowledge discovery approaches that can efficiently exploit this modern computing infrastructure. (<http://web.ornl.gov/sci/knowledgediscovery/CloudComputing/BDAC-SC14/>)

DISCS-2014: International Workshop on Data Intensive Scalable Computing Systems

9am-5:30pm

Room: 298-99

Philip C. Roth (Oak Ridge National Laboratory), Weikuan Yu (Auburn University), Yong Chen (Texas Tech University)

Existing high performance computing (HPC) systems are designed primarily for workloads requiring high rates of computation. However, the widening performance gap between processors and storage, and trends toward higher data intensity in scientific and engineering applications, suggest there is a need to rethink HPC system architectures, programming models, runtime systems, and tools with a focus on data intensive computing. This workshop builds on the momentum generated by its two predecessor workshops, providing a forum for researchers interested in HPC and data intensive computing to exchange ideas and discuss approaches for addressing Big Data challenges. The workshop includes a keynote address and presentation of peer-reviewed research papers, with ample opportunity for informal discussion throughout the day.

E2SC: Energy Efficient Supercomputing

9am-5:30pm

Room: 291

Darren J. Kerbyson (Pacific Northwest National Laboratory), Kirk Cameron (Virginia Tech), Adolfo Hoisie (Pacific Northwest National Laboratory), David Lowenthal (University of Arizona), Dimitris Nikolopoulos (Queen's University Belfast), Sudhakar Yalamanchili (Georgia Institute of Technology)

With exascale systems on the horizon, we have ushered in an era with power and energy consumption as the primary concerns for scalable computing. To achieve a viable exaflop high performance computing capability, revolutionary methods are required with a stronger integration among hardware features, system software and applications. Equally important are the capabilities for fine-grained spatial and temporal measurement and control to facilitate these layers for energy efficient computing across all layers. Current approaches for energy efficient computing rely heavily on power efficient hardware in isolation. However, it is pivotal for hardware to expose mechanisms for energy efficiency to optimize power and energy consumption for various workloads. At the same time, high fidelity measurement techniques, typically ignored in data-center level measurement, are of high importance for scalable and energy efficient inter-play in different layers of application, system software and hardware.

IA³ 2014: Fourth Workshop on Irregular Applications: Architectures and Algorithms

9am-5:30pm

Room: 273

Antonino Tumeo, John Feo (Pacific Northwest National Laboratory), Oreste Villa (NVIDIA Corporation)

Many data intensive applications are naturally irregular. They may present irregular data structures, control flow or communication. Current supercomputing systems are organized around components optimized for data locality and regular computation. Developing irregular applications on current machines demands a substantial effort, and often leads to poor performance. However, solving these applications efficiently is a key requirement for next generation systems. The solutions needed to address these challenges can only come by considering the problem from all perspectives: from micro-to system-architectures, from compilers to languages, from libraries to runtimes, from algorithm design to data characteristics. Only collaborative efforts among researchers with different expertise, including end users, domain experts, and computer scientists, could lead to significant breakthroughs. This workshop aims at bringing together scientists with all these different backgrounds to discuss, define and design methods and technologies for efficiently supporting irregular applications on current and future architectures.

Innovating the Network for Data Intensive Science

9am-12:30pm

Room: 274

Cees De Laat (University of Amsterdam), Jennifer M. Schopf, Martin Swamy (Indiana University)

Every year SCinet develops and implements the network for the SC conference. This network is state of the art, connecting many demonstrators of big data processing infrastructures at the highest line speeds and newest technologies available and demonstrates the newest functionality. The show floor network connects to many laboratories worldwide using Lambda connections and NREN networks. This workshop brings together the network researchers and innovators to bring up challenges and novel ideas that stretch SCinet even further. We invite papers that propose and discuss new and novel techniques regarding capacity and functionality of networks, its control and its architecture to be demonstrated at SC14.

Integrating Computational Science into the Curriculum: Models and Challenges

9am-12:30pm

Room: 293

Steven Gordon (Ohio Supercomputer Center)

Computational science has become the third path to discovery in science and engineering along with theory and experimentation. It is central to advancing research and is essential in making U.S. industry competitive in the face of international market competition. Yet, many universities have not integrated computational science into their curricula. This workshop will focus on the challenges of curriculum reform for computational science and the examples, alternative approaches, and existing and emerging resources that can be used to facilitate those changes. Participants will review competencies and models of computational science programs, learn about the demand for a workforce with computational modeling skills, and available resources and opportunities. Participants will analyze their current curriculum and expertise and devise a draft plan of action to advance computational science on their own campuses.

MTAGS 2014: 7th Workshop on Many-Task Computing on Clouds, Grids, and Supercomputers

9am-5:30pm

Room: 295

Ioan Raicu (Illinois Institute of Technology), Justin M. Wozniak (Argonne National Laboratory), Yong Zhao (University of Electronic Science and Technology of China)

The 7th workshop on Many-Task Computing on Clouds, Grids, and Supercomputers (MTAGS) will provide the scientific community a dedicated forum for presenting new research,

development, and deployment efforts of large-scale many-task computing (MTC) applications on large scale clusters, clouds, grids, and supercomputers. MTC, the theme of the workshop encompasses loosely coupled applications, which are generally composed of many-tasks to achieve some larger application goal. This workshop will cover challenges that can hamper efficiency and utilization in running applications on large-scale systems, such as local resource manager scalability and granularity, efficient utilization of raw hardware, parallel file-system contention and scalability, data management, I/O management, reliability at scale, and application scalability. We welcome paper submissions in theoretical, simulations, and systems topics with special consideration to papers addressing the intersection of petascale/exascale challenges with large-scale cloud computing. We invite the submission of original research work of 6 pages. For more information, see: <http://datasys.cs.iit.edu/events/MTAGS14>.

PDSW 2014: 9th Parallel Data Storage Workshop

9am-5:30pm

Room: 271-72

Dean Hildebrand (IBM Almaden Research Center), Garth Gibson (Carnegie Mellon University), Rob Ross (Argonne National Laboratory), Joan Digney (Carnegie Mellon University)

Peta- and exascale computing infrastructures make unprecedented demands on storage capacity, performance, concurrency, reliability, availability, and manageability. This one-day workshop focuses on the data storage and management problems and emerging solutions found in peta- and exascale scientific computing environments, with special attention to issues in which community collaboration can be crucial for problem identification, workload capture, solution interoperability, standards with community buy-in, and shared tools. Addressing storage media ranging from tape, HDD, and SSD, to new media like NVRAM, the workshop seeks contributions on relevant topics, including but not limited to performance and benchmarking, failure tolerance problems and solutions, APIs for high performance features, parallel file systems, high bandwidth storage architectures, support for high velocity or complex data, metadata intensive workloads, autonomies for HPC storage, virtualization for storage systems, archival storage advances, resource management innovations, and incorporation of emerging storage technologies.

Second Workshop on Sustainable Software for Science: Practice and Experiences (WSSSPE 2)

9am-5:30pm

Room: 275-76-77

Gabrielle D. Allen (University of Illinois at Urbana-Champaign), Daniel S. Katz (National Science Foundation), Karen Cranston (National Evolutionary Synthesis Center), Neil Chue Hong (Software Sustainability Institute, University of Edinburgh), Manish Parashar (Rutgers University), David Proctor (National Science Foundation), Matthew Turk (Columbia University), Colin C. Venters (University of Huddersfield), Nancy Wilkins-Diehr (San Diego Supercomputer Center, University of California, San Diego)

Progress in scientific research is dependent on the quality and accessibility of software at all levels and it is critical to address many new challenges related to the development, deployment, and maintenance of reusable software. In addition, it is essential that scientists, researchers, and students are able to learn and adopt a new set of software-related skills and methodologies. Established researchers are already acquiring some of these skills, and in particular a specialized class of software developers is emerging in academic environments as an integral and embedded part of successful research teams. Following a first workshop at SC13, WSSSPE2 will use reviewed short papers, keynote speakers, breakouts and panels to provide a forum for discussion of the challenges, including both positions and experiences. All material and discussions will be archived for continued discussion. The workshop is anticipated to lead to a special issue of the Journal of Open Research Software.

The 5th International Workshop on Performance Modeling, Benchmarking and Simulation of High Performance Computer Systems (PMBS14)

9am-5:30pm

Room: 283-84-85

Simon D. Hammond (Sandia National Laboratories), Stephen A. Jarvis (University of Warwick)

This workshop is concerned with the comparison of high-performance computing systems through performance modeling, benchmarking or through the use of tools such as simulators. We are particularly interested in research which reports the ability to measure tradeoffs in software/hardware co-design to improve sustained application performance. We are also keen to capture the assessment of future systems, for example through work that ensures continued application scalability through exascale systems.

The aim of this workshop is to bring together researchers, from industry and academia, concerned with the qualitative and quantitative evaluation and modeling of high-performance

computing systems. Authors are invited to submit novel research in areas of performance modeling, benchmarking and simulation. We welcome research that brings together theory and practice. We recognize that the coverage of the term 'performance' has broadened to include power and reliability, and that performance modeling is practiced through analytical methods and approaches based on software tools.

Ultravis '14: The 9th Workshop on Ultrascale Visualization

9am-5:30pm

Room: 294

Kwan-Liu Ma (University of California, Davis), Venkatram Vishwanath (Argonne National Laboratory), Hongfeng Yu (University of Nebraska-Lincoln)

The output from leading-edge scientific simulations and experiments is so voluminous and complex that advanced visualization techniques are necessary to interpret the calculated results. Even though visualization technology has progressed significantly in recent years, we are barely capable of exploiting petascale data to its full extent, and exascale datasets are on the horizon. This workshop aims at addressing this pressing issue by fostering communication between visualization researchers and the users of visualization. Attendees will be introduced to the latest and greatest research innovations in large-scale data visualization, and also learn how these innovations impact scientific supercomputing and discovery.

WHPCF'14: Seventh Workshop on High Performance Computational Finance

9am-5:30pm

Room: 292

Jose Moreira, David Daly (IBM), Matthew Dixon (University of San Francisco)

The purpose of this workshop is to bring together practitioners, researchers, vendors, and scholars from the complementary fields of computational finance and high performance computing, in order to promote an exchange of ideas, develop common benchmarks and methodologies, discuss future collaborations and develop new research directions. Financial companies increasingly rely on high performance computers to analyze high volumes of financial data, automatically execute trades, and manage risk.

Recent years have seen the dramatic increase in compute capabilities across a variety of parallel systems. The systems have also become more complex with trends towards heterogeneous systems consisting of general-purpose cores and acceleration devices. The workshop will enable the dissemination of recent advances and findings in the application of high performance computing to computational finance among

researchers, scholars, vendors and practitioners, and will encourage and highlight collaborations between these groups in addressing high performance computing research challenges.

WORKS 2014: 9th Workshop on Workflows in Support of Large-Scale Science

9am-5:30pm

Room: 297

Johan Montagnat (French National Center for Scientific Research), Ian Taylor (Cardiff University)

Data Intensive Workflows (a.k.a. scientific workflows) are routinely used in most scientific disciplines today, especially in the context of parallel and distributed computing. Workflows provide a systematic way of describing the analysis and rely on workflow management systems to execute the complex analyses on a variety of distributed resources. This workshop focuses on the many facets of data-intensive workflow management systems, ranging from job execution to service management and the coordination of data, service and job dependencies. The workshop therefore covers a broad range of issues in the scientific workflow lifecycle that include: data intensive workflows representation and enactment; designing workflow composition interfaces; workflow mapping techniques that may optimize the execution of the workflow; workflow enactment engines that need to deal with failures in the application and execution environment; and a number of computer science problems related to scientific workflows such as semantic technologies, compiler methods, fault detection and tolerance.

EduHPC: Workshop on Education for High Performance Computing

2pm-5:30pm

Room: 293

Sushil Prasad (Georgia State University), Almadena Chtchelkanova (National Science Foundation), Anshul Gupta (IBM Research), Arnold Rosenberg (Northeastern University), Alan Sussman (University of Maryland), Charles Weems (University of Massachusetts)

Parallel and Distributed Computing (PDC), especially its aspects pertaining to HPC, now permeates most computing activities. Certainly, it is no longer sufficient for even basic programmers to acquire only the traditional sequential programming skills. This workshop on state of art in high performance, parallel, and distributed computing education will comprise contributed as well as invited papers from academia, industry, and other educational and research institutes on topics pertaining to the teaching of PDC and HPC topics in the Computer Science and Engineering, Computational Science, and Domain Science and Engineering curriculum. The emphasis of the workshop will be on the undergraduate education,

although graduate education issues are also within scope, and target audience will include attendees among SC-14 Educators, academia, and industry. This effort is in coordination with NSF/TCPP curriculum initiative on parallel and distributed computing (www.cs.gsu.edu/~tcpp/curriculum/index.php). This workshop was the first education-related regular workshop held at SC-13.

NDM'14: Fourth International Workshop on Network-Aware Data Management

2pm-5:30pm

Room: 274

Mehmet Balman (VMware & LBNL), Surendra Byna (Lawrence Berkeley National Laboratory), Brian Tierney (Energy Sciences Network)

Data sharing and resource coordination among distributed teams are becoming significant challenges every passing year. Networking is one of the most crucial components in the overall system architecture of a data centric environment. Many of the current solutions both in industry and scientific domains depend on the underlying network infrastructure and its performance. There is a need for efficient use of the networking middleware to address increasing data and compute requirements. Main scope of this workshop is to promote new collaborations between data management and networking communities to evaluate emerging trends and current technological developments, and to discuss future design principles of network-aware data management. We will seek contribution from academia, government, and industry to address current research and development efforts in remote data access mechanisms, end-to-end resource coordination, network virtualization, analysis and management frameworks, practical experiences, data-center networking, and performance problems in high-bandwidth networks.

Monday, November 17

4th Workshop on Python for High Performance and Scientific Computing (PyHPC)

9am-5:30pm

Room: 291

Andreas Schreiber (German Aerospace Center), William Scullin (Argonne National Laboratory), Andy R. Terrel (Continuum Analytics)

Python is an established, general-purpose, high-level programming language with a large following in research and industry for applications in fields including computational fluid dynamics, finance, biomolecular simulation, artificial intelligence, statistics, data analysis, scientific visualization, and systems management. The use of Python in scientific, high performance parallel, big data, and distributed computing roles

has been on the rise with the community providing new and innovative solutions while preserving Python's famously clean syntax, low learning curve, portability, and ease of use.

The workshop will bring together researchers and practitioners from industry, academia, and the wider community using Python in all aspects of high performance and scientific computing. The goal is to present Python applications from mathematics, science, and engineering, to discuss general topics regarding the use of Python, and to share experience using Python in scientific computing education. For more information, see www.dlr.de/sc/pyhpc2014.

5th Annual Energy Efficient HPC Working Group Workshop

9am-5:30pm

Room: 271-72

Natalie Bates (Energy Efficient High Performance Computing Working Group), Stephen Poole (Oak Ridge National Laboratory), John Shalf (Lawrence Berkeley National Laboratory), Herbert Huber (Leibniz Supercomputing Center), Anna Maria Bailey (Lawrence Livermore National Laboratory), Susan Coghlan (Argonne National Laboratory), James Laros (Sandia National Laboratories), Josip Loncaric (Los Alamos National Laboratory), James Rogers (Oak Ridge National Laboratory), Wu Feng (Virginia Tech), Daniel Hackenberg (Technical University Dresden), Ladina Gilly (Swiss National Supercomputing Center), Dave Martinez (Sandia National Laboratories), William Tschudi (Lawrence Berkeley National Laboratory), Mariann Silveira (Lawrence Livermore National Laboratory), Thomas Durbin (University of Illinois at Urbana-Champaign), Kevin Regimbal (National Renewable Energy Laboratory), Francis Belot (French Alternative Energies and Atomic Energy Commission)

This annual workshop is organized by the Energy Efficient HPC Working Group (<http://eehpcwg.lbl.gov/>). It provides a strong blended focus that includes both the facilities and system perspectives; from architecture through design and implementation. The topics reflect the activities and interests of the EE HPC WG, which is a group with over 400 members from ~20 different countries. Speakers from SC13 included Chris Malone, Google, Dan Reed, University of Iowa and Jack Dongarra, University of Tennessee. There were also panel sessions covering all of the EE HPC WG Team activities. Panel topics from SC13 included lessons learned from commissioning liquid cooling building infrastructure, a methodology for improved quality power measurements for benchmarking and re-thinking the PUE metric. Dynamic speakers and interesting panel sessions characterized the SC13 workshop and can be expected for the SC14 workshop as well.

5th Workshop on Latest Advances in Scalable Algorithms for Large-Scale Systems (ScalA)

9am-5:30pm

Room: 283-84-85

Vassil Alexandrov (Barcelona Supercomputing Center), Al Geist, Christian Engelmann (Oak Ridge National Laboratory)

Novel scalable scientific algorithms are needed to enable key science applications to exploit the computational power of large-scale systems. This is especially true for the current tier of leading petascale machines and the road to exascale computing as HPC systems continue to scale up in compute node and processor core count. These extreme-scale systems require novel scientific algorithms to hide network and memory latency, have very high computation/communication overlap, have minimal communication, and have no synchronization points. Scientific algorithms for multi-petaflop and exaflop systems also need to be fault tolerant and fault resilient, since the probability of faults increases with scale. With the advent of heterogeneous compute nodes that employ standard processors and GPGPUs, scientific algorithms need to match these architectures to extract the most performance. Key science applications require novel mathematical models and system software that address the scalability and resilience challenges of current- and future-generation extreme-scale HPC systems.

Asian Technology Information Program (ATIP) Workshop on Japanese Research Toward Next-Generation Extreme Computing

9am-5:30pm

Room: 295

David Kahaner (Asian Technology Information Program), Satoshi Matsuoka (Tokyo Institute of Technology)

This workshop will include a significant set of presentations, posters, and panel discussions by Japanese researchers from universities, government laboratories, and industry. Participants will address topics including national exascale plans as well as the most significant hardware and software research. A key aspect of the proposed workshop will be the unique opportunity for members of the US research community to interact and have direct discussions with the top Japanese scientists who are participating. SC is the ideal venue for this workshop, because after the US, more SC participants come from Japan than from any other country. There are a multitude of exhibitor booths, research papers, and panels, etc. with Japanese content and Japanese researchers frequently win awards for best performance, greenest system, fastest networks, etc.

Co-HPC: Hardware-Software Co-Design for High Performance Computing

9am-5:30pm

Room: 273

Shirley Moore (University of Texas at El Paso), Richard Vuduc (Georgia Institute of Technology), Gregory Peterson (University of Tennessee, Knoxville), Theresa Windus (Iowa State University)

Hardware-software co-design involves the concurrent design of hardware and software components of complex computer systems, whereby application requirements influence architecture design and hardware constraints influence design of algorithms and software. Concurrent design of hardware and software has been used for the past two decades for embedded systems in automobiles, avionics, mobile devices, and other such products, to optimize for design constraints such as performance, power, and cost. HPC is facing a similar challenge as we move towards the exascale era, with the necessity of designing systems that run large-scale simulations with high performance while meeting cost and energy consumption constraints. This workshop will invite participation from researchers who are investigating the interrelationships between algorithms/applications, systems software, and hardware, and who are developing methodologies and tools for hardware-software co-design for HPC.

ExaMPI14: Exascale MPI 2014

9am-5:30pm

Room: 286-87

Stefano Markidis, Erwin Laure (KTH Royal Institute of Technology), Jesper Larsson Träff (Vienna University of Technology), Daniel Holmes, Mark Bull (Edinburgh Parallel Computing Center), William Gropp (University of Illinois at Urbana-Champaign), Masamichi Takagi (NEC Corporation), Lorna Smith, Mark Parsons (Edinburgh Parallel Computing Center)

The MPI design and its main implementations have proved surprisingly scalable. For this and many other reasons MPI is currently the de-facto standard for HPC systems and applications. However, there is a need for re-examination of the Message Passing (MP) model and for exploring new innovative and potentially disruptive concepts and algorithms, possibly to investigate other approaches than those taken by the MPI 3.0 standard.

The aim of workshop is to bring together researchers and developers to present and discuss innovative algorithms and concepts in the MP programming model and to create a forum for open and potentially controversial discussions on the future of MPI in the exascale era.

Possible workshop topics include innovative algorithms for collective operations, extensions to MPI, including data centric models such as active messages, scheduling/routing to avoid network congestion, “fault-tolerant” communication, interoperability of MP and PGAS models, and integration of task-parallel models in MPI.

Extreme-Scale Programming Tools

9am-5:30pm

Room: 297

Michael Gerndt (Technical University of Munich), Judit Gimenez (Barcelona Supercomputing Center), Martin Schulz (Lawrence Livermore National Laboratory), Felix Wolf (German Research School for Simulation Sciences), Brian J. N. Wylie (Juelich Supercomputing Centre)

Approaching exascale, architectural complexity and severe resource limitations with respect to power, memory and I/O make tools support in debugging and performance optimization more critical than ever before. However, the challenges mentioned above also apply to tools development and, in particular, raise the importance of topics such as automatic tuning and methodologies for exascale tools-aided application development. This workshop will serve as a forum for application, system, and tool developers to discuss the requirements for future exascale-enabled tools and the roadblocks that need to be addressed on the way. We also highly encourage application developers to share their experiences with using tools.

This workshop is the third in a series at SC conferences organized by the Virtual Institute - High Productivity Supercomputing (VI-HPS), an international initiative of HPC programming-tool builders. The event will also focus on the community-building process necessary to create an integrated tools-suite ready for an exascale software stack.

HPTCDL: First Workshop for High Performance Technical Computing in Dynamic Languages

9am-5:30pm

Room: 293

Jiahao Chen (Massachusetts Institute of Technology), Wade Shen (Massachusetts Institute of Technology and Lincoln Laboratories), Andy Terrel (Continuum Analytics)

Dynamic high-level languages such as Julia, Maple®, Mathematica®, MATLAB®, Octave, Python, R, and Scilab are rapidly gaining popularity with computational scientists and engineers, who often find these languages more productive for rapid prototyping of numerical simulation codes. However, writing legible yet performant code in dynamic languages remains challenging, which limits the scalability of code written in such languages, particularly when deployed on massively parallel architectures such as clusters, cloud servers, and

supercomputers. This workshop aims to bring together users, developers, and practitioners of dynamic technical computing languages, regardless of language, affiliation or discipline, to discuss topics of common interest. Examples of such topics include performance, software development, abstractions, composability and reusability, best practices for software engineering, and applications in the context of visualization, information retrieval and big data analytics.

LLVM-HPC: LLVM Compiler Infrastructure in HPC

9am-5:30pm

Room: 292

Hal Finkel, Jeff Hammond (Argonne National Laboratory)

LLVM, winner of the 2012 ACM Software System Award, has become an integral part of the software-development ecosystem for optimizing compilers, dynamic-language execution engines, source-code analysis and transformation tools, debuggers and linkers, and a whole host of programming-language and tool chain-related components. Now heavily used in both academia and industry, where it allows for rapid development of production-quality tools, LLVM is increasingly used in work targeted at high-performance computing. Research in and implementation of programming-language analysis, compilation, execution and profiling has clearly benefited from the availability of a high-quality, freely-available infrastructure on which to build. This workshop will focus on recent developments, from both academia and industry, which build on LLVM to advance the state of the art in high-performance computing.

VISTech 2014: Visualization Infrastructure and Systems Technology

9am-5:30pm

Room: 294

Kelly P. Gaither (Texas Advanced Computing Center, University of Texas at Austin), Jason Leigh (University of Hawaii at Manoa), Falko Kuester (University of California, San Diego), Eric A. Wernert (Indiana University), Aditi Majumder (University of California, Irvine), Karla P. Vega (Indiana University)

Human perception is centered on the ability to process information contained in visible light, and our visual interface is a tremendously powerful data processor. Every day we are inundated with staggering amounts of digital data. For many types of computational research, the field of visualization is the only viable means of extracting information and developing understanding from this data. Integrating our visual capacity with technological capabilities has tremendous potential for transformational science. We seek to explore the intersection between human perception and large-scale visual analysis through the study of visualization interfaces and interactive displays. This rich intersection includes: virtual reality systems, visualization through augmented reality, large scale visualiza-

tion systems, novel visualization interfaces, high-resolution interfaces, mobile displays, and visualization display middleware. The VISTech workshop will provide a space for experts in the large-scale visualization technology field and users to come together to discuss state-of-the art technologies for visualization and visualization laboratories.

WACCPD: First Workshop on Workshop on Accelerator Programming using Directives

9am-5:30pm

Room: 275-76-77

Sunita Chandrasekaran (University of Houston), Oscar Hernandez, Fernanda Foertter (Oak Ridge National Laboratory)

The nodes of many current HPC platforms are equipped with hardware accelerators that offer high performance with power benefits. In order to enable their use in scientific application codes without undue loss of programmer productivity, several recent efforts have been devoted to providing directive-based programming interfaces. These APIs promise application portability and a means to avoid low-level accelerator-specific programming. Many application developers prefer incremental ways to port codes to accelerator using directives without adding more complexity to their code. This workshop explores the use of these directive sets their implementations and experiences with their deployment in HPC applications. The workshop aims at bringing together the user and tools community to share their knowledge and experiences of using directives to program accelerators.

WOLFHPC14: Fourth International Workshop on Domain-Specific Languages and High-Level Frameworks for High Performance Computing

9am-5:30pm

Room: 298-99

Sriram Krishnamoorthy (Pacific Northwest National Laboratory), J. Ramanujam (Louisiana State University), P. Sadayappan (Ohio State University)

Multi-level heterogeneous parallelism and deep memory hierarchies in current and emerging computer systems makes their programming very difficult. Domain-specific languages (DSLs) and high-level frameworks (HLFs) provide convenient abstractions, shielding application developers from much of the complexity of explicit parallel programming in standard programming languages like C/C++/Fortran. However, achieving scalability and performance portability with DSLs and HLFs is a significant challenge. For example, very few high-level frameworks can make effective use of accelerators such as GPUs and FPGAs. This workshop seeks to bring together developers and users of DSLs and HLFs to identify challenges and discuss solution approaches for their effective implementation and use on massively parallel systems.

Friday, November 21

Advancing On-line HPC Learning

8:30am-12pm

Room: 286-87

Scott Lathrop (Shodor Education Foundation)

The goal of this workshop is to identify the challenges and opportunities for developing and delivering the infrastructure, content, teaching methods, certification and recognition mechanisms for providing high quality on-line programs. Another key goal is to produce publicly available reporting that will document the lessons learned and recommendations for advancing the development and delivery of high quality on-line programs.

DataCloud2014: 5th International Workshop on Data Intensive Computing in the Clouds

8:30am-12pm

Room: 292

Yong Zhao (University of Electronic Science and Technology of China), Ziming Zheng (HP Vertica), Wei Tang (Argonne National Laboratory)

Applications and experiments in all areas of science are becoming increasingly complex. Some applications generate data volumes reaching hundreds of terabytes and even petabytes. As scientific applications become more data intensive, the management of data resources and dataflow between the storage and compute resources is becoming the main bottleneck. Analyzing, visualizing, and disseminating these large data sets has become a major challenge and data intensive computing is now considered as the “fourth paradigm” in scientific discovery after theoretical, experimental, and computational science.

The 5th international workshop on Data-intensive Computing in the Clouds (DataCloud 2014) will provide the scientific community a dedicated forum for discussing new research, development, and deployment efforts in running data-intensive computing workloads on Cloud Computing infrastructures. The workshop will focus on the use of cloud-based technologies to meet the new data intensive scientific challenges that are not well served by the current supercomputers, grids or compute-intensive clouds.

GCE14: 9th Gateway Computing Environments Workshop

8:30am-12pm

Room: 273

Nancy Wilkins-Diehr (San Diego Supercomputer Center), Marlon E. Pierce (Indiana University), Suresh Marru (Indiana University)

Science today is increasingly digital and collaborative. The impact of high-end computing has exploded as new communities accelerate their research through science gateways

such as CIPRES and iPlant. Currently 40% of the NSF XSEDE program’s users come through science gateways. As datasets increase in size, communities increasingly use gateways for remote analysis. Software has scalable broader impact when researchers set up web interfaces to up-to-date codes running on high-end resources. Gateways increasingly connect varied elements of cyberinfrastructure - instruments, streaming sensor data, data stores and computing resources of all types. Online collaborative tools allow the sharing of both source data and subsequent analyses, speeding discovery. The important work of gateway development, however, is often done in an isolated, hobbyist environment. Leveraging knowledge about common tasks frees developers to focus on higher-level, grand-challenge functionality in their discipline. This workshop will feature case studies and an opportunity to share common experiences.

HUST14: First International Workshop on HPC User Support Too

8:30am-12pm

Room: 297

Ralph C. Bording (IVEC, University of Western Australia), Andy Georges (Ghent University), Kenneth Hoste (Ghent University)

Researchers pushing the boundaries of science and technology are an existential reason for supercomputing centers. To be productive, they heavily depend on HPC support teams, who in turn often struggle to adequately support the researchers.

Nevertheless, recent surveys have pointed out that there is an abundant lack of collaboration between HPC support teams all around the world, even though they are frequently facing very similar problems with respect to providing end users with the tools and services they require.

With this workshop we aim to bring together all parties involved, i.e. system administrators, user support team members, tool developers, policy makers and end users, to discuss these issues and bring forward the solutions they have come up with. As such, we want to provide a platform to present tools, share best practices and exchange ideas that help streamline HPC user support.

SEHPCSE14: Second International Workshop on Software Engineering for HPC in CSE

8:30am-12pm

Room: 283-84-85

Jeffrey Carver (University of Alabama), Neil Chue Hong (Software Sustainability Institute, University of Edinburgh), Selim Ciraci (Pacific Northwest National Laboratory)

Researchers are increasingly using high performance computing (HPC), including GPGPUs and computing clusters, for computational science & engineering (CSE) applications. Unfortunately, when developing HPC software, developers must solve reliability, availability, and maintainability problems in extreme

scales, understand domain specific constraints, deal with uncertainties inherent in scientific exploration, and develop algorithms that use computing resources efficiently. Software engineering (SE) researchers have developed tools and practices to support development tasks, including: validation and verification, design, requirements management and maintenance. HPC CSE software requires appropriately tailored SE tools/methods. The SE-HPCCSE workshop addresses this need by bringing together members of the SE and HPC CSE communities to share perspectives, present findings from research and practice, and generating an agenda to improve tools and practices for developing HPC CSE software. In the 2013 edition of this workshop, the discussion focused around a number of interesting topics, including: bit-by-bit vs. scientific validation and reproducibility.

VPA: First Workshop on Visual Performance Analysis

8:30am-12pm

Room: 298-99

Peer-Timo Bremer (Lawrence Livermore National Laboratory), Bernd Mohr (Forschungszentrum Juelich), Valerio Pascucci (University of Utah), Martin Schulz (Lawrence Livermore National Laboratory)

Over the last decades an incredible amount of resources has been devoted to building ever more powerful supercomputers. However, exploiting the full capabilities of these machines is becoming exponentially more difficult with each generation of hardware. To help understand and optimize the behavior of massively parallel simulations the performance analysis community has created a wide range of tools to collect performance data, such as flop counts or network traffic at the largest scale. However, this success has created new challenges, as the resulting data is too large and too complex to be analyzed easily. Therefore, new automatic analysis and visualization approaches must be developed to allow application developers to intuitively understand the multiple, interdependent effects that their algorithmic choices have on the final performance. This workshop intends to bring together researchers from performance analysis and visualization to discuss new approaches of combining both areas to analyze and optimize large-scale applications.

Women in HPC

8:30am-12pm

Room: 275-76-77

Toni Collis, Daniel Holmes, Lorna Smith, Alison Kennedy (Edinburgh Parallel Computing Centre)

Gender inequality is a problem across all scientific disciplines. Women are more likely to successfully complete tertiary education than men but less likely to become scientists. The multi-disciplinary background of HPC scientists should facilitate broad female engagement but fewer than 10% of participants

were female at two recent HPC conferences. Strong gender stereotyping of science negatively impacts female uptake and achievement in scientific education and employment. Removing gender stereotyping engages more women and gender-balanced groups have greater collective intelligence, which should benefit the HPC community.

This workshop will begin to address gender inequality in HPC by encouraging on-going participation of female researchers and providing an opportunity for female early career researchers to showcase their work and network with role-models and peers in an environment that reduces the male gender stereotype. Invited talks from leading female researchers will discuss their careers and a panel session will create a gender-equality action plan.

Workshop on Best Practices for HPC Training

8:30am-12pm

Room: 294

Fernanda Foertter (Oak Ridge National Laboratory), Rebecca Hartman-Baker (iVEC), Richard Gerber (Lawrence Berkeley National Laboratory), Nia Alexandrov (Barcelona Supercomputing Center), Barbara Chapman (University of Houston), Kjersten Fagnan (Lawrence Berkeley National Laboratory), Scott Lathrop (University of Illinois at Urbana-Champaign), Henry Neeman (University of Oklahoma), Maria-Ribera Sancho (BarcelonaTech), Robert Whitten (University of Tennessee, Knoxville)

HPC facilities face the challenge of serving a diverse user base with different skill levels and needs. Some users run precompiled applications, while others develop complex, highly optimized codes. Therefore, HPC training must include a variety of topics at different levels to cater to a range of skillsets. As centers worldwide install increasingly heterogeneous architectures, training will be even more important and in greater demand. A good training program can have many benefits: less time spent on rudimentary assistance, efficient utilization of resources, increased staff-user interaction, and training of the next generation of users. Unfortunately, the most successful training strategies at HPC facilities are not documented. This workshop aims to expose best practices for delivering HPC training. Topics will include: methods of delivery, development of curricula, optimizing duration, surveys and evaluations, metrics and determining success. Lastly, the workshop aims to develop collaborative connections between participating HPC centers.



Call for Participation

<http://sc15.supercomputing.org/>

The International Conference for High Performance Computing, Networking, Storage and Analysis

Conference Dates: Nov. 15 - 20, 2015 | Exhibition Dates: Nov. 16-19, 2015

SC15: HPC Transforms

High Performance Computing (HPC) is transforming our everyday lives, as well as our not-so-ordinary ones. From nanomaterials to jet aircrafts, from medical treatments to disaster preparedness, and even the way we wash our clothes; the HPC community has transformed the world in multifaceted ways.

For its 27th anniversary, the annual SC Conference will return to Austin, TX, a city that continues to develop new ways of engaging our senses and incubating technology of all types. SC15 will yet again provide a unique venue for spotlighting HPC and scientific applications and innovations from around the world.

SC15 will bring together the international supercomputing community—an unparalleled ensemble of scientists, engineers, researchers, students, educators, programmers, system administrators, developers, and managers—for an exceptional program of technical papers, informative tutorials, timely research posters, Birds-of-a-Feather (BOF) sessions, and much more. The SC15 Exhibition Hall will feature exhibits of the latest and greatest technologies from industry, academia and government research organizations; many of these technologies will be making their debut in Austin.

No conference is better poised to demonstrate how HPC can transform both the everyday and the incredible.



Austin Convention Center

Sponsors:



Association for
Computing Machinery

