

SC25-W: Workshops of the International Conference for High Performance Computing, Networking, Storage and Analysis

**St. Louis, Missouri, USA
16-21 November 2025**

Pages 1–689

IEEE Catalog Number: CFP25VZ4-POD
ISBN: 979-8-3195-0493-7



Copyright © 2025, ACM

All Rights Reserved

Copyright and Reprint Permissions:

Abstracting is permitted with credit to the source. For permission to reprint or republish any portion of this publication, please contact the copyright holder.

*** This is a print representation of what appears in an electronic format. Some format issues inherent in the e-media version may also appear in this print version.

IEEE Catalog Number:	CFP25VZ4-POD
ISBN (Print-On-Demand):	979-8-3195-0493-7
ISBN (Online):	979-8-4007-1871-7

Additional Copies of This Publication Are Available From:

Curran Associates, Inc
57 Morehouse Lane
Red Hook, NY 12571 USA

Phone: (845) 758-0400
Fax: (845) 758-2633
E-mail: curran@proceedings.com
Web: www.proceedings.com

TABLE OF CONTENTS

Parallel Data Object Creation: Scalable Metadata Management in Parallel I/O Library	1
<i>Youjia Li, Robert Latham, Robert Ross, Ankit Agrawal, Alok Choudhary, Wei-keng Liao</i>	
An Elastic Job Scheduler for HPC Applications on the Cloud	12
<i>Aditya Bhosale, Kavitha Chandrasekar, Laxmikant Kale, Sara Kokkila-Schumacher</i>	
Guiding Application Users Via Estimation of Computational Resources for Massively Parallel Chemistry Computations	24
<i>Tanzila Tabassum, Omer Subasi, Ajay Panyala, Epiya Ebiapia, Gerald Baumgartner, Erdal Mutlu</i> <i>P. Saday Sadayappan, Karol Kowalski</i>	
InferA: A Smart Assistant for Cosmological Ensemble Data	33
<i>Justin Z. Tam, Pascal Grosset, Divya Banesh, Nesar Ramachandra, Terece L. Turton, James Ahrens</i>	
Inverse Design for Generating Initial Conditions in Scientific Simulations	42
<i>Leslie Horace, Christin Whitton, Vanessa Job, William Jones, Nathan DeBardeleben</i>	
Applying Surrogate Modeling to Decouple Data Collection and Analysis from Simulation for Accelerated In-Situ Analysis	50
<i>Kewei Yan, Yonghong Yan</i>	
Compute4Biology: Taking Stock of High Performance Computing Needs for Foundation Models in Biological Sciences	58
<i>Pratik Dutta, Tirthankar Ghosal</i>	
FIRST: Federated Inference Resource Scheduling Toolkit for Scientific AI Model Access	65
<i>Aditya Tanikanti, Benoit Côté, Yanfei Guo, Le Chen, Nickolaus Saint, Ryan Chard, Ken Raffanetti</i> <i>Rajeev Thakur, Thomas Uram, Ian Foster, Michael E. Papka, Venkatram Vishwanath</i>	
ROSE: RADICAL Orchestrator for Surrogate Exploration	74
<i>Aymen Alsaadi, Tianle Wang, Andrew Park, Pradeep Bajracharya, Linwei Wang, Fanbo Sun, Sudip Seal</i> <i>Vikram Jadhao, Geoffrey Fox, Shantenu Jha</i>	
InferCT: An Efficient and Generalizable Framework to Enable 3D Machine Learning for Computed Tomography	84
<i>Austin Yunker, Weijian Zheng, Rajkumar Kettimuthu</i>	
LangChain-Parsl: Connect Large Language Model Agents to High Performance Computing Resource	91
<i>Heng Ma, Alexander Brace, Carlo Siebenschuh, Ian Foster, Arvind Ramanathan</i>	
A Query Engine for Scientific Data Exploration Using Theory, Simulation, and Artificial Intelligence Models	99
<i>Abhishek Dwaraki, Sreenivas. R. Sukumar, Christopher D. Rickett, Clarete Riana Crasta, Harumi Kuno</i> <i>Michael Neal, Pankaj Pandey, Ryan Yates, Karlon West, Amar Gopal Chittiboyina, Ikhlas A. Khan</i> <i>John L. Byrne, Sekwon Lee, David Emberson</i>	
Classification of Three-Dimensional Electron Diffraction Data with a Large Language Model	109
<i>Kazuyuki Yasuda, Masahito Kumagai, Masayuki Sato, Kazuhiko Komatsu, Hiroaki Kobayashi</i>	
Experience Deploying Containerized GenAI Services at an HPC Center	117
<i>Angel M Beltre, Jeff Ogden, Kevin Pedretti</i>	

Engine-Agnostic Model Hot-Swapping for Cost-Effective LLM Inference	127
<i>Radostin Stoyanov, Viktória Spišáková, Adrian Reber, Wesley Armour, Marcin Copik, Rodrigo Bruno</i>	
Seamless End-to-End Containerized HPC Environments	139
<i>Brandon Cook, Shane Canon, Adam Lavelly, Daniel Margala</i>	
Usability Evaluation of Cloud for HPC Applications	148
<i>Vanessa Sochat, Daniel Milroy, Abhik Sarkar, Aniruddha Marathe, Tapasya Patki</i>	
Evaluating HPK for Running Cloud-Native Workloads on Slurm Clusters	164
<i>Antony Chazapis, Lefteris Vassilakis, Giannis Petsis, Manolis Marazakis, Angelos Bilas</i>	
Towards Enabling Hostile Multi-Tenancy in Kubernetes	173
<i>Ali Kanso, Slava Oks, Mostafa Elzeiny, Gurpreet Virdi</i>	
Using Code Coverage to Assess Feature Gaps in MPI Correctness Tool Classification Tests	180
<i>Alexander Hück, Simon Schwitanski, Tim Jammer, Joachim Jenke, Yussur Mustafa Oraj, Christian Bischof</i>	
Coupling Static and Dynamic MPI Correctness Tools to Optimize Accuracy and Overhead	189
<i>Yussur Mustafa Oraj, Simon Schwitanski, Semih Burak, Christian Bischof, Matthias Müller</i>	
Data Race Detection Through Vibe Translation	199
<i>Jan Hueckelheim, Vimarsh Sathia, Siyuan Brant Qian</i>	
Differential Testing for Sequential to Parallel Transformations	208
<i>Jobayer Ahmmed, Quazi I. Mahmud, Junhyung Shim, Liyi Li, Ali Jannesari, Myra B. Cohen</i>	
Towards an Automated Workflow for Floating-Point Analysis of GPU Kernels	218
<i>Esteban Miguel Rangel, S. John Pennycook</i>	
LLM4FP: LLM-Based Program Generation for Triggering Floating-Point Inconsistencies Across Compilers	226
<i>Yutong Wang, Cindy Rubio-González</i>	
Exploring Reduced Precision for Deep Learning Activation Functions	236
<i>Epifanio Sarinana, Christoph Lauter, Shirley Moore</i>	
Extending MPI Correctness Benchmarking to the Fortran Language	245
<i>Yussur Mustafa Oraj, Alexander Hück, Christian Bischof</i>	
Design and Implementation of a Custom Hardware Accelerator for SZx Compression in Chipyard	250
<i>Connor Bohannon, Kazutomo Yohsui, Sheng Di, Franck Cappello, Antonino Miceli</i>	
ASCRIBE-XR: Extended Reality for Visualization of Scientific Images	260
<i>Ronald Pandolfi, Julian Todd, Jeffrey J Donatelli, Daniela Ushizima</i>	
Characterizing the Performance of Parallel Data-Compression Algorithms Across Compilers and GPUs	270
<i>Brandon Alexander Burtchell, Martin Burtscher</i>	
Integrating Distributed SQL Query Engines with Object-Based Computational Storage	280
<i>Junhyun Ryu, Soon Hwang, Junhyeok Park, Seonghoon Ahn, JeoungAhn Park, Jeongjin Lee, Jinna Yang, Soonyeal Yang, Jungki Noh, Qing Zheng, Woosuk Chung, Hoshik Kim, Youngjae Kim</i>	
On the Compressibility of Floating-Point Data in Posit and IEEE-754 Representation	290
<i>Andrew Rodriguez, Martin Burtscher</i>	

Building n-Dimensional Trees for Resolution-Based Progressive Compression	297
<i>Brandon Alexander Burtchell, Martin Burtscher</i>	
Lightweight CNN-Based Artifact Reduction for Scientific Error-bounded Lossy Compression	304
<i>Zizhe Jian, Pu Jiao, Bohan Zhang, Sheng Di, Xin Liang, Guanpeng Li, Huangliang Dai, Zizhong Chen</i> <i>Franck Cappello</i>	
Benchmarking Cutting-Edge Scientific Error-Bounded Lossy Compressors on Correlation-Based Rate-Distortion	314
<i>Ziwei Qiu, Jinyang Liu, Kai Zhao, Robert Underwood, Sheng Di</i>	
FZModules: A Heterogeneous Computing Framework for Customizable Scientific Data Compression Pipelines	322
<i>Skyler Ruiter, Jiannan Tian, Fengguang Song</i>	
Compression Error Sensitivity Analysis for Different Experts in MoE Model Inference	329
<i>Songkai Ma, Zhaorui Zhang, Sheng Di, Benben Liu, Xiaodong Yu, Xiaoyi Lu, Dan Wang</i>	
Evaluating Accuracy and Performance Tradeoffs in GPU Accelerated Single Cell RNA-seq Analysis	339
<i>Cory Gardner, Seyun Jeong, Oam Khatavkar, Aiden Moon, Qinglei Cao, Tae-Hyuk Ahn</i>	
Hybrid GPU Programming Education with Python and C++: Preferences, Performance and Common Python Pitfalls	349
<i>Lena Oden</i>	
An Interactive Agentic HPC Tutor for Lesson Planning, Teaching, and Assessment	357
<i>Erik Pautsch, Mengjiao Han, Joseph A. Insley, Janet Knowles, Victor A. Mateevitsi, Silvio Rizzi</i> <i>George K. Thiruvathukal</i>	
Teaching Task-Based Parallel Programming with a Runtime Systems-Aware Perspective	366
<i>Vivek Kumar</i>	
GPU Programming for AI Workflow Development on AWS SageMaker: An Instructional Approach	374
<i>Sriram Srinivasan, Hamdan Alabsi, Rand Obeidat, Nithisha Ponnala, Azene Zenebe</i>	
The Cost of Teaching Operational ML	383
<i>Fraida Fund, Kate Keahey, Cody Hammock, Marc Richardson, Mark Powers, Michael Sherman</i>	
A Model for Teaching Machine Learning, Deep Learning, and Research Computing to Domain Scientists on HPC Resources	391
<i>William J. Allen, Kelsey M. Beavers, Erik Ferlanti, Lorenzo Concia, Joshua Urrutia, Ernesto A. B. F. Lima</i> <i>John M. Fonner, Felix Zuo, H. E. Duplechin Seymour, Ari B. Kahn, Joe Stubbs, Anagha Jamthe</i> <i>Stephanie N. Baker, Tabish Khan, James P. Carson</i>	
Peachy Parallel Assignments (EduHPC 2025)	399
<i>Clara Almeida, Elizabeth Shoop, Diego García-Álvarez, Arturo Gonzalez-Escribano</i> <i>David Guerrero-Pantoja, Cameron Maloney, Maria Pantoja, Silvio Rizzi, David P. Bunde</i>	
From Soil to Software: Experience from a STEM Workshop on Smart Plant Care and Teachable Machines	406
<i>Anita Esmaeilian, Kishwar Ahmed</i>	
“Offloading” Undergraduate Research to the Graphics Processing Unit for Acceleration	415
<i>Bryant Wyatt, Mason Bane</i>	

MPPi - Type Safe C++ Datatypes for MPI	421
<i>Mike Söhner, Christoph Niethammer</i>	
Large-Message All-to-All Communication at Frontier Scale	431
<i>James Buford White</i>	
MPI Collectives with Programmable Smart Switches	438
<i>Thomas Erbesdobler, Amir Raoofy, Ehab Saleh, Josef Weidendorfer</i>	
Scaling All-to-All Operations Across Emerging Many-Core Supercomputers	449
<i>Shannon Gayle Kinkad, Jackson Wesley, Whit Schonbein, David DeBonis, Matthew Dosanjh Amanda Bienz</i>	
On the Integration of Lightweight Tasks with MPI Using the C++26 Std::Execution 'Senders' API	459
<i>John Biddiscombe, Mikael Simberg, Auriane Reverdell, Raffaele Solcà, Alberto Invernizzi, Rocco Meli Joseph Schuchart</i>	
Enabling Unstructured Sparse Fine-Tuning and Inference for Foundation Models on Wafer-Scale Engine	470
<i>Haoyu Zheng, Yifan Zeng, Linghao Song, Murali Emani, Wenqian Dong</i>	
WAGES: Workload-Aware GPU Sharing System for Energy-Efficient Serverless LLM Serving	478
<i>Tianyu Wang, Gourav Rattihalli, Aditya Dhakal, Xulong Tang, Dejan Milojicic</i>	
OmniFed: A Modular Framework for Configurable Federated Learning from Edge to HPC	486
<i>Sahil Tyagi, Andrei Cozma, Olivera Kotevska, Feiyi Wang</i>	
Enhancing ChatPORT with CUDA-to-SYCL Kernel Translation Capability	494
<i>Zheming Jin, Swaroop Pophale, Keita Teranishi</i>	
Evaluation of Test-Time Compute Constraints on Safety and Skill Large Reasoning Models	504
<i>Adarsha Balaji, Le Chen, Rajeev Thakur, Franck Cappello, Sandeep Madireddy</i>	
Batch Tiling on Attention: Efficient Mixture of Experts Training on Wafer-Scale Processors	510
<i>Daria Soboleva, Etienne Goffinet, Hui Zeng, Sangamesh Ragate, Elif Albuz, Natalia Vassilieva</i>	
Automated MCQA Benchmarking at Scale: Evaluating Reasoning Traces as Retrieval Sources for Domain Adaptation of Small Language Models	515
<i>Ozan Gokdemir, Neil Getty, Robert Underwood, Sandeep Madireddy, Franck Cappello, Arvind Ramanathan Ian T. Foster, Rick L. Stevens</i>	
Gridmind: LLMs-Powered Agents for Power System Analysis and Operations	523
<i>Hongwei Jin, Kibaek Kim, Jonghwan Kwon</i>	
Frameworks for Large Language Model Serving in HPC Environments	532
<i>Rohan Marwaha, Qinren Zhou, Kastan Day, Asmita Dabholkar, Volodymyr Kindratenko</i>	
Exploring Distributed Vector Databases Performance on HPC Platforms: A Study with Qdrant	538
<i>Seth Ockerman, Amal Gueroudji, Song Young Oh, Robert Underwood, Nicholas Chia, Kyle Chard Robert Ross, Shivaram Venkataraman</i>	
EQSIM Agent: A Conversational AI for Interactive Exploration of Large-Scale Earthquake Simulation Data	545
<i>Houjun Tang, David McCallen</i>	

Beyond End-to-End: Understanding the Limits of LLMs in Scientific Problem Solving	551
<i>Youyuan Liu, Sheng Di, Neil Getty, Tanwi Mallick, Robert Underwood, Sian Jin</i>	
BioR5: A Three-Layer Architecture for Biological Reasoning in Scientific AI	557
<i>Peng Ding, Thomas Brettin, Rick Stevens</i>	
LABMATE: Language Model Based Multi-Agent System to Accelerate Catalysis Experiments	565
<i>Anurag Acharya, Anshu Kiran Sharma, Derek Parker, Timothy Vega, Rizwan A. Ashraf, Natalie M. Isenberg Jan Strube, Robert Rallo</i>	
Programmer Productivity and Performance on AMD's AI Engines: Offloading Fortran Intrinsics via MLIR a Case-Study	574
<i>Nick Brown, Gabriel Rodriguez-Canal</i>	
Connected-Component Labeling Using HIs for High-Energy Particle Physics Instruments	582
<i>Nick Song, Marion Sudvarg, Roger Chamberlain</i>	
Aurora Acceptance: A Collaborative Exascale Test Harness	591
<i>Brian Homerding, Longfei Gao, Ben Lenard, Eric Pershey, Kevin Harms, Doug Waldron, Susan Coghlan Bill Allcock, Ti Leggett, Balazs Gerofi, Tom Musta</i>	
Experiences Integrating Database Support Into the OLCF Test Harness	602
<i>Nick Hagerty, Daniel Dietz</i>	
Testing and Benchmarking Emerging Supercomputers Via the MFC Flow Solver	609
<i>Benjamin Wilfong, Anand Radhakrishnan, Henry Le Berre, Tanush Prathi, Stephen Abbott Spencer H. Bryngelson</i>	
Seeking Cost-Optimal Infrastructure Size for Distributed Filesystems: A Ceph Case Study	618
<i>Niccolo Tosato, Isac Pasianotto, Ruggero Lot, Stefano Cozzini</i>	
A Modular, Responsive, and Accessible HPC Dashboard Built Upon Open OnDemand	628
<i>Richie Tan, Guangzhen Jin</i>	
Open Composer: A Web-Based Application for Generating and Managing Batch Jobs on HPC Clusters	637
<i>Masahiro Nakao, Keiji Yamamoto</i>	
Is it an HPC Workflow Assistant? Is it a Framework? It's Drona Workflow Engine	645
<i>Andrii Kryvenko, Duy Pham, Marinus Pennings, Honggao Liu</i>	
Generating Frequently Asked Questions from Technical Support Tickets Using Large Language Models	655
<i>Christina Joslin, David Burns, Fnu Ashish, Elham Sarbijan</i>	
AskHPC: A ChatBot for High Performance Computing User Support	667
<i>Akhilesh Bondapalli, Huihuo Zheng, Oluwaseun T. Ajayi, Murat Keceli, Haritha Siddabathuni Som Taylor Childers, Lisa Childers, Yasaman Ghadar, Michael Papka, Venkatram Vishwanath, Rong Ge</i>	
ModuLair: Streamlining Python Virtual Environment Management for HPC	680
<i>Surada Suwansathit, Ananya Adiki, Gabriel Floreslovo, Marinus Pennings, Honggao Liu</i>	
Dori: User Centered HPC for Data Intensive Computing	690
<i>Georg Rath</i>	
eIM: GPU-Accelerated Efficient Influence Maximization in Large-Scale Social Networks	697
<i>Jacob Doney, Xin Huang, Chul-Ho Lee</i>	

Generating Permutations at Scale	706
<i>Oded Green, Joe Eaton, Alok Tripathy, Corey Nolet, Justin Luitjens</i>	
An Optimized Generalized Multi-Color Point Implicit Solver for Intel GPUs Using OneAPI ESIMD	715
<i>Joseph Wassell, Mohammad Zubair, Aaron Walden, Gabriel Nastac, Eric Nielsen, Timothée Ewart</i>	
Profiling Application-Specific Properties of Irregular Graph Algorithms on GPUs	724
<i>Bishal Sharma, Martin Burtscher</i>	
Performance-Portable Symbolic Factorization Through Common Graph Operations	733
<i>Oguz Selvitopi, Xiaoye S. Li, Aydin Buluc</i>	
Comparing Graph Algorithm Styles on NVIDIA and Amd GPUs	742
<i>Avery VanAusdal, Martin Burtscher</i>	
Benchmarking the Cerebras Wafer Scale Engine-2 Architecture	747
<i>Takaaki Miyajima, Leon Fukuoka</i>	
How Effective is Matrix Reordering for Improving Performance of Sparse Matrix-Vector Multiplication?	752
<i>Omid Asudeh, Sina Mahdipour Saravani, Fabrice Rastello, Gerald Sabin, Ponnuswamy Sadayappan</i>	
Rapid Quantum Network Simulation Design with a Path to Scalable Execution	757
<i>Aaron Welch, Joel Dawson, Mariam Kiran</i>	
To Stream or Not to Stream: Towards a Quantitative Model for Remote HPC Processing Decisions	763
<i>Flavio Castro, Weijian Zheng, Joaquin Chung, Ian Foster, Raj Kettimuthu</i>	
From Path-Aware to Application-Aware Source Routing Using Traffic Classes	770
<i>Shashwitha Puttaswamy, Mariam Kiran</i>	
Learning to Schedule: A Supervised Learning Framework for Network-Aware Scheduling of Data-Intensive Workloads	777
<i>Sankalpa Timilsina, Susmit Shannigrahi</i>	
LLM-Based Optimization Algorithm Selection for High-Performance Networks Orchestration	783
<i>Anestis Dalgkitsis, Cyril Hsu, Chrysa Papagianni, Paola Grosso, Cees de Laat</i>	
Optimizing Network Resilience Using Domain-Specific Hardware Accelerator for Dynamic Programming	789
<i>Ali Mazloum, Sergio Elizalde, Samia Choueiri, Elie Kfoury, Jose Gomez, Ali AlSabeh, Jorge Crichigno</i>	
Complex Parsing for In-Network Acceleration of High-Energy Physics Experiments	795
<i>Bjoern Sagstad, Nishanth Shyamkumar, Sophia Chen, Wesley Robert Ketchum, Roland Sipos, James Kowalkowski, Michael Wang, Nik Sultana</i>	
SENSE in Practice: Quantifying the End-to-End Benefits of Intent-Based Bandwidth Reservation for Exascale Science Workflows	804
<i>Inder Monga, Mazahir Hussain, Justas Balcas, Aashay Arora, Diego Davila, Buseung Cho, Xi Yang, Cees De Laat</i>	
Scaling LLM Training Using RDMA Over Converged Ethernet	815
<i>Alex Batlle Casellas, Adrián Pérez Diéguez, Aleix Torres-Camps, Harris Teague, Arnau Padres Masdemont, Jordi Ros-Giralt</i>	

Network Replay and Consistency Across Testbeds	826
<i>Alexander Wolosewicz, Vinod Yegneswaran, Ashish Gehani, Nik Sultana</i>	
eCounter: Inline Per-IP Network Monitoring at Millisecond Resolution via eBPF	837
<i>Xinxin Mei, Jie Chen, Amitoj Singh, Ilya Baldin, David Lawrence</i>	
Implementing Network-level QoS at HPC Datacenters to Enable Distributed Scientific Workflows	848
<i>Anna Giannakou, Jonathan Skone, Vinay Sawal, Ronal Kumar, Stephen Simms, Nicholas Wright Lavanya Ramakrishnan</i>	
StreamHub: High-performance Managed SciStream as a Service	858
<i>Seena VazifeDunn, Flavio Castro, Rajkumar Kettimuthu, Ian Foster, Kyle Chard</i>	
Modular Architecture for High-Performance and Low Overhead Data Transfers	868
<i>Rasman Muhtasim Swargo, Engin Arslan, Md Arifuzzaman</i>	
From Edge to HPC: Investigating Cross-Facility Data Streaming Architectures	878
<i>Anjus George, Michael Brim, Christopher Zimmer, David Rogers, Sarp Oral, Zach Mayes</i>	
SmartNIC Data Exchange Framework	889
<i>Zackary Savoie, Anthony Sicoie, Ryan Eric Grant</i>	
MPI Communication Performance on AMD MI300A: Microbenchmarks and Applications	897
<i>Goutham Kalikrishna Reddy Kuncham, Siyuan Zhang, Shoaib Mohammad, Chen-Chun Chen Dhabaleswar K. Panda</i>	
In-Transit Data Transport Strategies for Coupled AI-Simulation Workflow Patterns	905
<i>Harikrishna Tummalapalli, Riccardo Balin, Christine Simpson, Andrew Park, Aymen Alsaadi Andrew E Shao, Wesley Brewer, Shantenu Jha</i>	
From Exploration to Explanation: ML-Driven Causal Discovery for Datacenter Reliability at Scale	917
<i>Pavana Prakash, Rolando P. Hong Enriquez, Sergey Serebryakov, David Grant, Wesley Brewer Dejan Milojicic</i>	
A Data-Size Adaptive Approach to I/O of Poorly Load Balanced in Situ Data Extracts	923
<i>Caitlin Ross, Scott Wittenburg, Greg Eisenhauer, Corey Wetterer-Nelson</i>	
Lessons Learned: Template-Heavy C++ in Production HPC Runtime Systems	929
<i>Arne Hendricks</i>	
ASaP: Automatic Software Prefetching for Sparse Tensor Computations in MLIR	937
<i>Konstantinos Sotiropoulos, Jonas Skeppstedt, Per Stenström</i>	
Implementing OpenMP Offload Support in the AMD Next Generation Fortran Compiler	948
<i>Dominik Adamski, Sergio Afonso, Akash Banerjee, Pranav Bhandarkar, Kareem Ergawy, Andrew Gozillon Michael Klemm, Jan Leyonberg, Dan Palermo</i>	
CIRE: LLVM Analysis for Floating-Point Rounding Error Affected by Precision and Optimizations	959
<i>Tanmay Tirpankar, Cayden Lund, Ganesh Gopalakrishnan</i>	
Lowering and Runtime Support for Fortran's Multi-Image Parallel Features Using LLVM Flang, PRIF, and Caffeine	968
<i>Dan Bonachea, Katherine Rasmussen, Damian Rouson, Jean-Didier Paillex, Etienne Renault Brad Richardson</i>	

Scabbard: LLVM Instrumentation-Aided Race Checking in CPU/GPU Unified Memory for AMD GPUs	977
<i>Andrew Osterhout, Ignacio Laguna, Ganesh Gopalakrishnan</i>	
OpenSHMEM MLIR: A Dialect for Compile-Time Optimization of One-Sided Communications	986
<i>Michael Beebe, Benjamin Michalowicz, Andrew McNamara, Yash Kumar, Dhabaleswar K. Panda Yong Chen, Wendy Poole, Steve Poole</i>	
Distributed Sparse Tensor Computations in MLIR	999
<i>Miheer Vaidya, Shreya Singh, Devanshu Mantri, Michael Shannon Eydenberg, Brian Michael Kelley Sivasankaran Rajamanickam, Atanas Rountev, P. Sadayappan</i>	
An MLIR Pipeline for Offloading Fortran to FPGAs via OpenMP	1007
<i>Gabriel Rodriguez-Canal, David Katz, Nick Brown</i>	
Umpire: Portable Memory Management for High-Performance Computing Applications	1015
<i>Kristi Belcher, David Beckingsale</i>	
The MALL is Open: Exploring Shared Caches and Latency in AMD CDNA™ 3 GPUs	1021
<i>Andrew Tee, Nicholas Curtis, Noah Wolfe, Daniel Wong</i>	
CaRDS: Compiler-Aided Remote Data Structures	1028
<i>Brian Tauro, Ian Dougherty, Kyle Hale</i>	
Hardware-Software Co-Design of Iterative Filter-Update Numerical Methods Using Processing-in-Memory	1037
<i>Eric Tang, Tianyun Zhang, William Bradford, Farzana Ahmed Siddique, James C. Hoe, Kevin Skadron Franz Franchetti</i>	
Performance Portable Batched Linear Algebra Kernels for Transport Sweeps Using Kokkos	1044
<i>Gabriel Suau, Thierry Gautier, Ansar Calloo, Remi Baron, Romain Le Tellier</i>	
Extending RAJA Parallel Programming Abstractions with Just-In-Time Optimization	1056
<i>John Bowen, Konstantinos Parasyris, David Beckingsale, Tal Ben-Nun, Thomas Stitt, Giorgis Georgakoudis</i>	
Energy-Aware Performance Portability with OpenMP Dynamic Variants	1069
<i>Ayman Bourramouss el Maach, Adrian Munera, Sara Royuela</i>	
Preserving CUDA Syntax for SYCL Portability: A Thin C++ Abstraction without Kernel Migration	1082
<i>Esteban Miguel Rangel, Humza Qureshi</i>	
Roofline Analysis of Tightly-Coupled CPU-GPU Superchips: A Study on MI300A and GH200	1091
<i>Oscar Antepara, Leonid Oliker, Samuel Williams</i>	
LAMMPS-KOKKOS: Performance Portable Molecular Dynamics Across Exascale Architectures	1103
<i>Anders Johansson, Evan Weinberg, Christian Trott, Megan McCarthy, Stan Moore</i>	
Development of a Performance Portable Distributed FFT Interface on Top of the Kokkos Ecosystem	1119
<i>Yuuichi Asahi, Trévis Morvany, Thomas Padioleau, Julien Bigot</i>	
Comparing Distributed-Memory Programming Frameworks with Radix Sort	1129
<i>Michael P. Ferguson, Ryan D. Friesse, Shreyas Khandekar, Matt Drozt</i>	

Slicing is All You Need: Towards A Universal One-Sided Algorithm for Distributed Matrix Multiplication	1142
<i>Benjamin Brock, Renato Golin</i>	
DiOMP-Offloading: Toward Portable Distributed Heterogeneous OpenMP	1155
<i>Baodi Shan, Mauricio Araya-Polo, Barbara Chapman</i>	
Weak Scaling of NVSHMEM Applied To Hashed Distributed Structured Data	1168
<i>Andrew Davis, Hans Johansen, Xinfeng Gao, Stephen Guzik</i>	
Enhancing HPX with FleCSI: Automatic Detection of Implicit Task Dependencies	1180
<i>S. Davis Herring, Maxim Moraru, Scott Pakin, Julien Loiseau, Richard Berger, Philipp V. F. Edelmann Ben Bergen</i>	
Stackless vs. Stackful Coroutines: A Comparative Study for RDMA-Based Asynchronous Many-Task (AMT) Runtimes	1191
<i>Mia Reitz, Jonas Posner</i>	
KDRSolvers: Scalable, Flexible, Task-Oriented Krylov Solvers	1201
<i>David Kai Zhang, Rohan Yadav, Alex Aiken, Fredrik Kjolstad, Sean Treichler</i>	
LLMTailor: A Layer-Wise Tailoring Tool for Efficient Checkpointing of Large Language Models	1216
<i>Minqui Sun, Xin Huang, Luanzheng Guo, Nathan R. Tallent, Kento Sato, Dong Dai</i>	
SlimIO: Lightweight I/O Path Design for Write Isolation in FDP-backed In-Memory Databases	1225
<i>Sangyun Lee, Sungjin Byeon, Soon Hwang, Jaewan Park, Joo-Young Hwang, Junyoung Han Javier González, Awais Khan, Youngjae Kim</i>	
RL4Sys: A Lightweight System-Driven RL Framework for Drop-in Integration in System Optimization	1235
<i>Jiaxin Dong, Md. Hasanur Rashid, Helen Xu, Dong Dai</i>	
Secure in-Storage Execution of VTK Workloads on Modern Parallel NFS Data Servers	1244
<i>Qing Zheng, Brian Atkinson, Jason Lee, Daoce Wang, Gary Grider</i>	
Modelling Load Imbalance in Shared Memory Multicore Systems	1253
<i>Johannes Langguth, James Trotter, Xing Cai</i>	
A Peak Performance Model for All-to-all on Hierarchical Systems and its Applications	1262
<i>Rohini Uma-Vaideswaran, Joshua Romero, Daniel L. Dotson, David Appelhans, P. K. Yeung</i>	
Beyond Guess and Check: Quantifying the Fidelity of Proxy Applications	1272
<i>Si Chen, Simon Garcia de Gonzalo, Omar Aaziz, Jeanine Cook, Avani Wildani</i>	
Implications of Full-System Modeling for Superconducting Architectures	1286
<i>Kunal Pai, Mahyar Samani, Anusheel Nand, Jason Lowe-Power</i>	
PerfAnalyzer: Revealing Performance Trends Using Version Oriented Visual Analysis of Scientific Software	1293
<i>Sayef Azad Sakin, James Ahrens</i>	
Experiences of Porting Structured and Unstructured Stencil Applications to FPGA Using SYCL	1298
<i>Zadok Storkey, Steven Wright, Ian Gray</i>	

MoE-Inference-Bench: Performance Evaluation of Mixture of Expert Large Language and Vision Models	1304
<i>Krishna Teja Chitty-Venkata, Sylvia Howland, Golar Azar, Daria Soboleva, Natalia Vassilieva Siddhisanket Raskar, Murali Emani, Venkatram Vishwanath</i>	
Pretraining LLMs at Scale: Tuning Strategies and Performance Portability	1314
<i>Adrián Pérez Diéguez, Álex Batlle Casellas, Aleix Torres-Camps, Harris Teague, Jordi Ros-Giralt</i>	
Characterizing the Impact of GPU Power Management on an Exascale System	1326
<i>Mariana Toledo Costa, Antigoni Georgiadou, Tames B. White, Bruno Villasenor Alvarez, Jordà Polo Woong Shin, Philippe Olivier, Bronson Messer, Arthur Francisco Lorenzon</i>	
A GPU FFT Wrapper to Co-optimize Floating-Point Precision and Library Selection via Predictive Error Modeling	1336
<i>Julius Lehner, Eishi Arima, Martin Schulz</i>	
ILAN: The Interference- and Locality-Aware NUMA Scheduler	1346
<i>Edvin Mellberg, Axel Carlsson, Jing Chen, Miquel Pericàs</i>	
On the Performance and Scalability of Cloud Supercomputers: Insights from Eagle and Reindeer	1356
<i>Amirreza Rastegari, Prabhat Ram, Michael F. Ringenburt</i>	
A Cache Interaction Graph for Data Locality Optimization	1366
<i>Chao Jin, David Abramson, Mo'ath Qutaish, Mark Endrei, Bronis R. de Supinski</i>	
MT4G: A Tool for Reliable Auto-Discovery of NVIDIA and AMD GPU Compute and Memory Topologies	1376
<i>Stepan Vanecek, Manuel Walter Mußbacher, Dominik Größler, Urvij Saroliya, Martin Schulz</i>	
Fantastic Hardware Counters and How to Find Them: Automating the Detection of Noise-Resilient Performance Counters in HPC	1389
<i>Ahmad Tarraf, Alexander GeiB, Lukas Fuchs, Felix Wolf</i>	
Extending THAPI with CXI Hardware Counter Sampling for High Resolution NIC Telemetry	1403
<i>Nathan Nichols, Thomas Applencourt</i>	
Scalable, High-Fidelity Monitoring of Application Communication Patterns in Vernier	1413
<i>Jered Dominguez-Trujillo, Derek Schafer, Riley Shipley, Ryan Marshall, Nicholas Bacon, Maxim Moraru Galen Shipman, Anthony Skjellum, Patrick Bridges</i>	
PEAK: Cost-Adaptive Profiling in a Heartbeat	1426
<i>Yuheng Chen, Junjie Li, Chun-Yaung Lu, Yinzhi Wang</i>	
Fast on-Demand Memory Mapping for Shared Memory and Disaggregated Systems	1436
<i>Yuang Yan, Ryan Grant</i>	
Modeling the Carbon Footprint of HPC: The Top 500 and EasyC	1445
<i>Varsha Rao, Andrew A. Chien</i>	
TEGRA - Scaling Up Graph Processing with Disaggregated Computing	1454
<i>William Shaddix, Mahyar Samani, Jason Lowe-Power, Venkatesh Akella</i>	
An RDMA-First Object Storage System with SmartNIC Offload	1461
<i>Yu Zhu, Aditya Dhakal, Pedro Bruel, Gourav Rattihalli, Yunming Xiao, Johann Lombardi, Dejan Milojicic</i>	

Doceph: DPU-Offloaded Messaging in Ceph for Reduced Host CPU Utilization	1469
<i>Kyuli Park, Sungmin Yoon, Farid Talibli, Sungyong Park, Jae-Hyuck Kwak, Kimoon Jeong, Awais Khan Youngjae Kim</i>	
IzhiRISC-V - a RISC-V-Based Processor with Custom ISA Extension for Spiking Neuron Networks Processing with Izhikevich Neurons	1478
<i>Wiktor Jan Szczerek, Artur Podobas</i>	
RISC-V Vectorization Coverage for HPC: A TSVN-Based Analysis	1487
<i>Hung-Ming Lai, Pei-Hung Lin, Maya Gokhale, Ivy Peng, Hiren Patel, Jenq-Kuen Lee</i>	
A RISC-V Vector Extension for Multi-word Arithmetic	1495
<i>Yunhao Lan, Larry Tang, Naifeng Zhang, Youngjin Eum, James Hoe, Franz Franchetti</i>	
Enabling the Syscall_Intercept Library for RISC-V	1505
<i>Petar Andrić, Ramon Nou, Aaron Call, Guillem Senabre</i>	
Is RISC-V Ready for High Performance Computing? An Evaluation of the Sophon SG2044	1514
<i>Nick Brown</i>	
Bridging Simulation and Silicon: A Study of RISC-V Hardware and FireSim Simulation	1523
<i>Atanu Barai, Kamalavasan Kamalakkannan, Patrick Diehl, Maxim Moraru, Jered Dominguez-Trujillo Howard Pritchard, Nandakishore Santhi, Farzad Fatollahi-Fard, Galen Shipman</i>	
Simulating Hybrid Analog + RISC-V Systems for HPC Applications	1534
<i>Cameron Durbin, Jacob Flores, Thomas Weatherly, Ben Feinberg</i>	
Accelerating Gravitational N-Body Simulations Using the RISC-V-Based Tenstorrent Wormhole	1540
<i>Jenny Lynn Almerol, Elisabetta Boella, Mario Spera, Daniele Gregori</i>	
Assessing a RISC-V Accelerator for Cross-Section Lookup in Chipyard	1547
<i>Andrew Ledbetter, Kazutomo Yoshii, John Tramm</i>	
Dyninst on the RISC-V: Binary Instrumentation in Support of Performance, Debugging, and Other Tools	1554
<i>Cheng-Hsun Angus He, Ronak Chauhan, James A. Kupsch, Hsuan-Heng Wu, Barton P. Miller</i>	
Extending the C++ Execution Control Library to Support Dynamic Parallel Runtime Systems	1562
<i>Ian Henriksen, Jan Ciesko, Stephen L. Olivier</i>	
Assessing Page Reclamation Mechanisms for Linux	1573
<i>Shaochang Liu, Jie Ren</i>	
Reproducible Performance Evaluation of OpenMP and SYCL Workloads Under Noise Injection	1581
<i>Christoffer Persson, Mathias Pretot, Minyu Cui, Miquel Pericas</i>	
Numerical Properties and Scalability of s-Step Preconditioned Conjugate Gradient Methods	1590
<i>Viktoria Mayer, Wilfried N. Gansterer</i>	
Efficient Embedding Initialization via Dominant Eigenvector Projections	1601
<i>Quentin Petit, Chong Li, Nahid Emad, Jack Dongarra</i>	
Fast Linear Solvers via AI-Tuned Markov Chain Monte Carlo-Based Matrix Inversion	1611
<i>Anton Lebedev, Won Kyung Lee, Soumyadip Ghosh, Olha I. Yaman, Vassilis Kalantzis, Yingdong Lu Tomasz Nowicki, Shashanka Ubaru, Lior Horesh, Vassil Alexandrov</i>	

A High Performance GPU CountSketch Implementation and Its Application to Multisketching and Least Squares Problems	1619
<i>Andrew James Higgins, Erik Boman, Ichitaro Yamazaki</i>	
Post-Variational Quantum Neural Networks on a Hybrid HPC-QC System	1627
<i>Maxence Vandromme, Miwako Tsuji</i>	
High-Performance and Power-Efficient Emulation of Matrix Multiplication Using INT8 Matrix Engines	1635
<i>Yuki Uchino, Katsuhisa Ozaki, Toshiyuki Imamura</i>	
An HPC-Inspired Blueprint for a Technology-Agnostic Quantum Middle Layer	1643
<i>Stefano Markidis, Gilbert Netzer, Luca Pennati, Ivy Peng</i>	
Towards a User-Centric HPC-QC Environment	1651
<i>Aleksander Wennersteen, Matthieu Moreau, Aurelien Nober, Mourad Beji</i>	
Orchestrating Quantum-HPC Workflows with Distributed Quantum Circuit Cutting	1660
<i>Mar Tejedor, Berta Casas, Javier Conejero, Alba Cervera-Lierta, Rosa M. Badia</i>	
Towards Supporting QIR: Steps for Adopting the Quantum Intermediate Representation	1669
<i>Yannick Stade, Lukas Burgholzer, Robert Wille</i>	
A Practical Quantum Solver for Multidimensional Partial Differential Equations	1678
<i>Manu Chaudhary, Kareem El-Araby, Alvir Nobel, Ishraq Islam, Manish Singh, Sunday Ogundele Kieran Egan, Sneha Thomas, Vincent Vordtriede, Devon Bontrager, Serom Kim, Esam El-Araby</i>	
Securing HDF5 Plugins with Digital Signatures	1688
<i>Glenn Song, Michael Scot Breitenfeld, Suren Byna</i>	
CASSE: Targeted Threat Modeling for Data Management Libraries	1696
<i>Keegan Sanchez, Suren Byna, Zhiqiang Lin, David Mattson</i>	
Dynamic Factor Graphs for Attack Preemption	1705
<i>Phuong M Cao, Zbigniew T Kalbarczyk, Ravishankar K Iyer</i>	
Evaluating Trusted Execution Environment Performance for Genome Sequence Alignment: An AMD SEV Case Study	1713
<i>Robert Keßler, Lech Nieroda, Simon Volpert, Moritz Gräf, Viktor Achter, Laslo Hunhold, Stefan Wesner</i>	
HPC Digital Twins for Evaluating Scheduling Policies, Incentive Structures and their Impact on Power and Cooling	1721
<i>Matthias Maiterth, Wesley H. Brewer, Jaya S. Kuruvella, Arunavo Dey, Tanzima Z. Islam, Rashadul Kabir Kevin Menear, Dmitry Duplyakin, Tapasya Patki, Terry Jones, Feiyi Wang</i>	
Run-Time Energy-Efficiency Optimization for AI and HPC Workloads	1732
<i>Gabriel Hautreux, Abdoulaye Gamatié, Gilles Sassatelli</i>	
Improving Supercomputer Usage with Aging Awareness	1742
<i>Robin Boezennec, Fanny Dufossé, Guillaume Pallez, Alix Tremodeux</i>	
Optimizing Microgrid Composition for Sustainable Data Centers	1752
<i>Julius Irion, Philipp Wiesner, Jonathan Bader, Odej Kao</i>	
Energy-Aware HPC Scheduling with LLM-Based Power Prediction	1759
<i>Kevin Menear, Alex Wilkinson, Tim Dykes, Utz-Uwe Haus, Dmitry Duplyakin</i>	

Bridging the Gap: User-Centric Energy Monitoring for Policy-Driven Application Optimization in HPC Data Centers	1769
<i>Woong Shin, Karl W. Schulz, Arthur F. Lorenzon, Matthias Maiterth, Bruno Villasenor Alvarez, Jordà Polo Aditya Kashi, Hao Lu, Nicholson Koukpaizan, Antigoni Georgiadou, Matthew Norman, Wael Elwasif Michael Matheson, Feiyi Wang, Nicholas Frontiere, Sarp Oral, Thomas Beck, Bronson Messer</i>	
Molten Chloride Small Modular Reactor Performance Characteristics for Data Center Operation	1779
<i>Matthew Anderson, Daniel Yankura, Matthew Sgambati, Mauricio Tano Retamales</i>	
EMLIO: Minimizing I/O Latency and Energy Consumption for Large-Scale AI Training	1784
<i>Md Hasibul Jamil, MD S. Q. Zulkar Nine, Tevfik Kosar</i>	
EAS-Sim: A Framework and its Methodology for the Co-Design of Multi-Objective, Energy-Aware Schedulers for AI Clusters	1794
<i>Roblex Nana Tchakoute, Claude Tadonki</i>	
Porting a Fortran plasma simulation to Exascale on AMD GPUS using both OpenMP and Kokkos	1804
<i>Etienne Malaboeuf, Kévin Obrejan, Mathieu Peybernes, Julien Bigot, Emily Bourne, Virginie Grandgirard</i>	
Bridging FPGA and GPU Over PCIe: A Low-Latency Communication Path Using AVX-512	1821
<i>Michele Martinelli, Carlotta Chiarini, Andrea Biagioni, Paolo Cretaro, Ottorino Frezza Francesca Lo Cicero, Alessandro Lonardo, Pierpaolo Perticaroli, Francesco Simula, Luca Pontisso Cristian Rossi, Piero Vicini</i>	
Towards Efficient Load Balancing BFS on GPUs: One Code for AMD, Intel & Nvidia	1830
<i>Kaan Olgu, Tobias Kenter, Jose Nunez-Yanez, Simon McIntosh-Smith, Tom Deakin</i>	
Scalable Neural Network Training: Distributed Data-Parallel Approaches	1841
<i>Fernando Vázquez-Novoa, Pedro López, José Flich, Rosa M. Badia</i>	
Reduction-Aware Directive-Based Programming via Multi-Dimensional Homomorphisms	1853
<i>Richard Schulze, Sergei Gorbach, Ari Rasch</i>	
A Study of Performance Portability of Low-Bit Fused Matrix-Vector Multiplication Kernels in SYCL	1867
<i>Zheming Jin</i>	
Physical System Study on Balancing Interactive and Batch Job Performance through Oversubscribing Scheduling	1875
<i>Shohei Minami, Toshio Endo, Akihiro Nomura, Hiroki Ohtsuji, Jun Kato, Masahiro Miwa, Eiji Yoshida</i>	
Implementing Support for Interactive and AI Workloads in a Traditional HPC Environment	1884
<i>Jay McGlothlin, Christopher Carothers</i>	
Modeling and Optimizing Real-Time Telescope Interaction for Multi-Wavelength Observation of Gamma-Ray Bursts	1889
<i>Ye Htet, Marion Sudvar, Honghao Yang, Jeremy Buhler, Roger Chamberlain, James Buckley</i>	
Adapting Classic Scheduling Heuristics for Online Execution Under Uncertainty	1897
<i>Jason Chamorro, Gabriel Twigg-Ho, Jared Coleman, Tainã Coleman, Bhaskar Krishnamachari Mohammadali Khodabandehlou</i>	
A Workflow for Error Analysis for Drug Response Prediction via Statistical Standardization and Distribution Analysis	1909
<i>Jake Gwinn, Justin M. Wozniak, Rajeev Jain, Yitan Zhu, Alex Partin, Thomas Bretin, Rick Stevens</i>	

Bridging Speed and Optimality in Job Scheduling: A Hybrid Ant Colony Optimization Approach for Distributed Systems	1918
<i>Hongwei Jin, Pawel Zuk, Krishnan Raghavan, Prachi Jadhav, Aiden Hamade, Ewa Deelman Prasanna Balaprakash</i>	
CAMEO: A Co-design Architecture for Multi-objective Energy System Optimization	1929
<i>Rounak Meyur, Sam Donald, Tonya Martin, Thiagarajan Ramachandran, Sumit Purohit</i>	
DAGonStore: Reliable Data Management for Workflows on the Computing Continuum with DynoStore and DAGonStar	1941
<i>Dante D. Sanchez-Gallegos, J. L. Gonzalez-Compean, Jesus Carretero, Raffaele Montella</i>	
Do Large Language Models Speak Scientific Workflows?	1953
<i>Orcun Yildiz, Tom Peterka</i>	
Evaluating the Efficacy of LLM-Based Reasoning for Multiobjective HPC Job Scheduling	1962
<i>Prachi Jadhav, Hongwei Jin, Ewa Deelman, Prasanna Balaprakash</i>	
Integrating and Characterizing HPC Task Runtime Systems for hybrid AI-HPC workloads	1973
<i>Andre Merzky, Mikhail Titov, Matteo Turilli, Shantenu Jha</i>	
Overcoming Dynamic I/O Boundaries: A Double-Sided Streaming Methodology with dispel4py and CAPIO	1985
<i>Marco Edoardo Santimaria, Rosa Filgueira, Doriana Medić, Iacopo Colonnelli, Marco Aldinucci</i>	
RESILIO: A Scalable and Composable Architecture for Tomographic Reconstruction Workflows	1997
<i>Amal Gueroudji, Matthieu Dorier, Philip Carns, Parth Patel, Tekin Bicer, Robert Latham, Robert Ross Kyle Chard, Ian Foster</i>	
State Machine Orchestration of an HPC Workflow in Cloud	2009
<i>Vanessa Sochat, Laic Pottier, Daniel Milroy</i>	
xGFabric: Coupling Sensor Networks and HPC Facilities with Private 5G Wireless Networks for Real-Time Digital Agriculture	2021
<i>Liubov Kurafeeva, Alan Subedi, Ryan Hartung, Michael Fay, Avhishek Biswas, Shantenu Jha, Ozgur Kilic Chandra Krintz, Andre Merzky, Douglas Thain, Mehmet Vuran, Rich Wolski</i>	
Accelerating Advanced Light Source Science Through Multi-Facility HPC Workflows	2032
<i>David Abramov, Samuel Welborn, Ryan Chard, Kuldeep Chawla, Xiaoya Chong, Elizabeth Clark Bjoern Enders, Alexander Hexemer, Jason Jed, Wiebke Koepp, Harinarayan Krishnan, Seij De Leon Dilworth Parkinson, David Perlmutter, Raja Vyshnavi Sriramoju, Thomas Uram, Lee Lisheng Yang Dylan McReynolds</i>	
Streaming X-ray Detector Data to Remote Facilities Using EJFAT	2040
<i>Siniša Veseli, John Hammonds, Steven Henke, Madeline Miller, Hannah Parraga, Ilya Baldin Derek Howard, Yatish Kumar, Nicholas Schwarz</i>	
Adapting Scientific Streaming Inference Workflows for a Deterministic Tensor Processing Unit	2051
<i>Samantha Fowler, Senthil Gnanasekaran, Kazutomo Yoshii, Antonino Miceli, Tao Zhou, Nicholas Contini</i>	
X-Ray Ptychography at the Edge: Towards Real-Time Feedback for High-Speed Nanoimaging	2058
<i>Zirui Gao, Seher Karakuzu, Dmitri Gavrilov, Daniel Allan, Adam Thompson, Denis Leshchev, Hanfei Yan</i>	