

1. Record Nr.	UNISA996466292803316
Titolo	High Performance Computing [[electronic resource]] : ISC High Performance 2019 International Workshops, Frankfurt, Germany, June 16-20, 2019, Revised Selected Papers / / edited by Michèle Weiland, Guido Juckeland, Sadaf Alam, Heike Jagode
Pubbl/distr/stampa	Cham : , : Springer International Publishing : , : Imprint : Springer, , 2019
ISBN	3-030-34356-1
Edizione	[1st ed. 2019.]
Descrizione fisica	1 online resource (XXV, 659 p. 402 illus., 239 illus. in color.)
Collana	Theoretical Computer Science and General Issues, , 2512-2029 ; ; 11887
Disciplina	004.3
Soggetti	Computer engineering Computer networks Software engineering Computers Computer Engineering and Networks Software Engineering Computer Hardware Computing Milieux
Lingua di pubblicazione	Inglese
Formato	Materiale a stampa
Livello bibliografico	Monografia
Nota di contenuto	Intro -- Preface -- Organization -- Short Papers -- Preface to the First International Workshop on Legacy Software Refactoring for Performance -- P ³ MA Workshop 2019 -- 4th International Workshop on In Situ Visualization (WOIV'19) -- Contents -- On the Use of Kernel Bypass Mechanisms for High-Performance Inter-container Communications -- 1 Introduction -- 2 Overview of Compared Solutions -- 3 Experimental Results -- 4 Related Work -- 5 Conclusions and Future Work -- References -- Continuous-Action Reinforcement Learning for Memory Allocation in Virtualized Servers -- 1 Introduction -- 2 Background -- 2.1 Memory Management in Virtualized Nodes -- 2.2 Reinforcement Learning: Markov Decision Process -- 3 CAVMem: Algorithm for Virtualized Memory Management -- 3.1 Decentralized Strategy for

Memory Management -- 3.2 Formulating the Problem as an MDP -- 4
 Experimental Framework -- 5 Results for Evaluation -- 5.1 Results for
 Scenario 1 -- 5.2 Results for Scenario 2 -- 5.3 Results for Scenario 3
 -- 5.4 Discussion -- 6 Related Work -- 7 Conclusions and Future Work
 -- References -- Container Orchestration on HPC Clusters -- 1
 Introduction -- 2 Related Work -- 3 Background -- 3.1 Kubernetes --
 3.2 Kubernetes Deployment -- 4 Implementation -- 4.1 General
 Approach -- 4.2 Kubernetes Cluster Deployment -- 4.3 HPC Worker
 Node Software Prerequisites -- 4.4 Networking -- 4.5 GE Worker Setup
 and Tear down -- 4.6 Kubernetes Cluster Configuration -- 5 Evaluation
 -- 6 Discussion -- 7 Conclusion and Future Work -- References --
 Data Pallets: Containerizing Storage for Reproducibility and Traceability
 -- 1 Introduction -- 2 Related Work -- 3 Design -- 3.1 Design and
 Implementation Challenges -- 3.2 Design and Implementation Details
 -- 3.3 Integration with Sandia Analysis Workbench (SAW) -- 4
 Measurements -- 4.1 Time Overheads -- 4.2 Space Overheads -- 4.3
 Discussion.
 5 Integration with Sandia Analysis Workbench -- 6 Conclusions and
 Future Work -- References -- Sarus: Highly Scalable Docker Containers
 for HPC Systems -- 1 Introduction -- 2 Related Work -- 3 Sarus -- 3.1
 Sarus Architecture -- 3.2 Container Creation -- 4 Extending Sarus with
 OCI Hooks -- 4.1 Native MPICH-Based MPI Support (H1) -- 4.2 NVIDIA
 GPU Support (H2) -- 4.3 SSH Connection Within Containers (H3) -- 4.4
 Slurm Scheduler Synchronization (H4) -- 5 Performance Evaluation --
 5.1 Scientific Applications -- 6 Conclusions -- References --
 Singularity GPU Containers Execution on HPC Cluster -- 1 Introduction
 -- 2 Singularity GPU Containers Building and Running -- 3 Benchmark
 -- 3.1 Systems Description -- 3.2 Test Case 1: Containerized
 Tensorflow Execution on GALILEO Versus Official Tensorflow
 Performance Data -- 3.3 Test Case 2: Containerized Versus Bare Metal
 Execution on GALILEO -- 4 Conclusion -- References -- A Multitenant
 Container Platform with OKD, Harbor Registry and ELK -- 1
 Introduction -- 2 Past -- 2.1 Background -- 2.2 Challenges -- 3
 Present -- 3.1 Evaluation of Container Orchestration Frameworks --
 3.2 Observability: Logging and OKD -- 3.3 Observability: Monitoring
 and OKD -- 4 Future -- 4.1 Monitoring -- 4.2 Container Policy and
 OKD -- 4.3 Gitops gitops and OKD -- 4.4 Continuous Delivery in OKD
 -- 4.5 OKD in the Cloud -- 5 Conclusion -- References -- Enabling
 GPU-Enhanced Computer Vision and Machine Learning Research Using
 Containers -- 1 Introduction -- 2 Defining the Base Container -- 2.1
 System Setup: Ubuntu, CUDA, Docker, Nvidia-Docker -- 2.2 Docker
 and Container Runtime -- 2.3 TensorFlow -- 2.4 OpenCV -- 2.5
 Cuda_tensorflow_opencv -- 3 Using the Base Container -- 3.1 Testing
 Code from a Bash Terminal -- 3.2 Integrating Darknet and Yolo V3
 Python Bindings -- 4 Conclusion -- References.
 Software and Hardware Co-design for Low-Power HPC Platforms -- 1
 Introduction -- 2 Network Interface Primitives -- 3 HPC Prototype -- 4
 User-Level Communication Library -- 5 MPI Implementation over the
 Proposed Architecture -- 6 Conclusions and Future Work -- References
 -- Modernizing Titan2D, a Parallel AMR Geophysical Flow Code to
 Support Multiple Rheologies and Extendability -- 1 Introduction -- 2
 Titan2D and Benchmark Problem -- 3 Refactoring Strategies -- 3.1
 Adopting a Python Interface -- 3.2 Merging Multiple Forks -- 3.3
 Changing Data Layout to for Modern CPU Architectures -- 3.4 Efficient
 Indexing for Elements/Nodes Addressing -- 3.5 Introducing OpenMP
 and Hybrid OpenMP/MPI Parallelization -- 4 Performance Improvement
 Evaluation -- 5 Conclusions and Future Plans -- References --
 Asynchronous AMR on Multi-GPUs -- 1 Introduction -- 2 Execution on

Heterogeneous Architectures -- 2.1 Data Model and CPU-GPU Communication -- 2.2 Scheduling on Heterogeneous Architectures -- 2.3 API -- 2.4 Multi-GPU Support -- 3 Evaluation -- 4 Conclusions -- References -- Batch Solution of Small PDEs with the OPS DSL -- 1 Introduction -- 2 The OPS DSL -- 3 Batching Support in OPS -- 3.1 Extending the Abstraction -- 3.2 Execution Schedule Transformation -- 3.3 Data Layout Transformation -- 3.4 Alternating Direction Implicit Solver -- 4 Evaluation -- 4.1 The Application -- 4.2 Experimental Set-Up -- 4.3 Results -- 5 Conclusions -- References -- Scalable Parallelization of Stencils Using MODA -- 1 Introduction -- 2 Related Work -- 3 Methodology -- 3.1 MODA and User-Defined Indices -- 3.2 Using GGDML Indices -- 3.3 Communication Identification -- 4 Evaluation -- 4.1 Test Application -- 4.2 Test System -- 4.3 Experiments -- 5 Summary -- References -- Comparing High Performance Computing Accelerator Programming Models -- 1 Introduction -- 2 Motivation -- 3 Related Work. 4 Analysis -- 5 Discussion -- 5.1 BT Benchmark -- 5.2 SP Benchmark -- 5.3 LBM Benchmark -- 5.4 LBDC Benchmark -- 6 Conclusion -- References -- Tracking User-Perceived I/O Slowdown via Probing -- 1 Introduction -- 2 Related Work -- 3 Methodology -- 3.1 Probing -- 3.2 Data Reduction Using Statistics -- 3.3 Computing the Slowdown -- 4 Evaluation -- 4.1 Test Systems -- 4.2 Probing Tool -- 4.3 Timeseries of Individual Measurements -- 4.4 Host Variability -- 4.5 Understanding Application Behavior - The IO-500 -- 4.6 Long-Period -- 4.7 Slowdown -- 5 Conclusion -- References -- A Quantitative Approach to Architecting All-Flash Lustre File Systems -- 1 Introduction -- 2 Methods -- 3 File System Capacity -- 4 Drive Endurance -- 5 Metadata Configuration -- 5.1 MDT Capacity Required by DOM -- 5.2 MDT Capacity Required for Inodes -- 5.3 Overall MDT Capacity -- 6 Conclusion -- References -- MBWU: Benefit Quantification for Data Access Function Offloading -- 1 Introduction -- 2 The MBWU-Based Methodology -- 2.1 Background -- 2.2 What Is MBWU -- 2.3 How to Measure MBWU(s) -- 2.4 Evaluation Prototype -- 3 Evaluation -- 3.1 Infrastructure -- 3.2 Test Setup and Results -- 4 Related Work -- 5 Conclusion -- References -- Footprinting Parallel I/O - Machine Learning to Classify Application's I/O Behavior -- 1 Introduction -- 2 Related Work -- 3 DKRZ Monitoring -- 3.1 Metrics -- 4 Methodology -- 5 Test Data -- 5.1 Data Preparation -- 6 Evaluation -- 6.1 I/O Behavior Classification -- 6.2 Footprinting -- 7 Manual Identification of I/O Intensive Jobs -- 8 Summary and Conclusion -- References -- Adventures in NoSQL for Metadata Management -- 1 Introduction -- 2 Related Work -- 3 Metadata Model -- 3.1 Basic Metadata -- 3.2 Custom Metadata -- 4 Design -- 4.1 What Has the Right Features to Be Worth Testing? -- 4.2 What Is It Going to Take to Get It All Working at All?. 4.3 Can We Make Our Queries Work with Any Performance? -- 4.4 Battle Scars and Lessons for Our Next Battle Against Scale Out Computing Tools -- 5 Evaluation -- 5.1 Insert Time -- 5.2 Query Time -- 6 Conclusion and Future Work -- References -- Towards High Performance Data Analytics for Climate Change -- 1 Introduction -- 2 Main Challenges -- 3 The Ophidia Project -- 3.1 Multi-dimensional Storage Model -- 3.2 Array-Based Primitives and Parallel Operators -- 4 Benchmark and Experimental Results -- 4.1 Benchmark Definition -- 4.2 Test Environment -- 4.3 Experimental Results and Discussion -- 5 Related Work -- 6 Conclusions -- References -- An Architecture for High Performance Computing and Data Systems Using Byte-Addressable Persistent Memory -- 1 Introduction -- 2 Persistent Memory -- 2.1 Data Access -- 2.2 B-APM Modes of Operation -- 2.3

Non-volatile Memory Software Ecosystem -- 3 Opportunities for Exploiting B-APM for Computational Simulations and Data Analytics -- 3.1 Potential Caveats -- 4 Systemware Architecture -- 4.1 Job Scheduler -- 4.2 Data Scheduler -- 5 Performance Evaluation -- 6 Related Work -- 7 Summary -- References -- Mediating Data Center Storage Diversity in HPC Applications with FAODEL -- 1 Introduction -- 2 FAODEL Background -- 2.1 Kelpie -- 2.2 I/O Management (IOM) Modules -- 3 Mediating Storage Using Kelpie Object Naming -- 3.1 Kelpie Architectural Considerations -- 3.2 Annotating the Kelpie Namespace -- 3.3 Service-Initiated Mediation -- 3.4 Performance Considerations -- 4 Related Work -- 5 Conclusion -- References -- Predicting File Lifetimes with Machine Learning -- 1 Introduction -- 2 Specifying the Problem and Building the Models -- 2.1 Problem Specification -- 2.2 Dataset -- 2.3 Data Preprocessing -- 2.4 Models -- 3 Results -- 3.1 Evaluation Methodology -- 3.2 Training Times and Model Sizes -- 3.3 Accuracy. 3.4 Error and Accuracy Distribution.

Sommario/riassunto

This book constitutes the refereed post-conference proceedings of 13 workshops held at the 34th International ISC High Performance 2019 Conference, in Frankfurt, Germany, in June 2019: HPC I/O in the Data Center (HPC-IODC), Workshop on Performance & Scalability of Storage Systems (WOPSSS), Workshop on Performance & Scalability of Storage Systems (WOPSSS), 13th Workshop on Virtualization in High-Performance Cloud Computing (VHPC '18), 3rd International Workshop on In Situ Visualization: Introduction and Applications, ExaComm: Fourth International Workshop on Communication Architectures for HPC, Big Data, Deep Learning and Clouds at Extreme Scale, International Workshop on OpenPOWER for HPC (IWOPH18), IXPUG Workshop: Many-core Computing on Intel, Processors: Applications, Performance and Best-Practice Solutions, Workshop on Sustainable Ultrascale Computing Systems, Approximate and Transprecision Computing on Emerging Technologies (ATCET), First Workshop on the Convergence of Large Scale Simulation and Artificial Intelligence, 3rd Workshop for Open Source Supercomputing (OpenSuCo), First Workshop on Interactive High-Performance Computing, Workshop on Performance Portable Programming Models for Accelerators (P³MA). The 48 full papers included in this volume were carefully reviewed and selected. They cover all aspects of research, development, and application of large-scale, high performance experimental and commercial systems. Topics include HPC computer architecture and hardware; programming models, system software, and applications; solutions for heterogeneity, reliability, power efficiency of systems; virtualization and containerized environments; big data and cloud computing; and artificial intelligence.
