

IEEE/ACM
2025 INTERNATIONAL
CONFERENCE ON
COMPUTER-AIDED
DESIGN

ICCAD

44th Edition

2025 INTERNATIONAL CONFERENCE ON COMPUTER-AIDED DESIGN CONFERENCE PROGRAM

SPONSORS AND ORGANIZERS



IEEE



Please visit our website for more information!

[2025.ICCAD.COM](https://2025.iccad.com)

Table of Contents

Executive Committee	3
Venue Map & Information.....	4
Welcome Message.....	7
Awards	9
BPA Candidates (Front End)	11
BPA Candidates (Back End)	13
Plenary Speakers.....	14
Tutorial Sessions	18
Workshops.....	20
Panel Session	26
CEDA Luncheon.....	27
Workforce Development	28
Social Functions	29
Sponsors & Exhibitors	30
Program at a Glance.....	32
Technical Program – Oct. 26	37
Technical Program – Oct. 27	39
Technical Program – Oct. 28	67
Technical Program – Oct. 29	98
Technical Program – Oct. 30	126

Executive Committee

General Chair

Robert Wille, Technical University of Munich, Munich Quantum Software Company, and Software Competence Center Hagenberg GmbH

Past Chair

Jinjun Xiong, University at Buffalo, USA

Program Chair

Deming Chen, University of Illinois, USA

Vice Program Chair

Ismail S. K. Bustany, AMD, USA

Tutorial & Special Session Chair

Tsung-Yi Ho, The Chinese University of Hong Kong, China

Workshop Chair

Ron Duncan, Synopsys

CEDA Representative

Jiang Hu, Texas A&M University, USA

ACM SIGDA Representative

Wanli Chang, Hunan University, China

Asian Representative

Wei Zhang, The Hong Kong University of Science and Technology, China

European Representative

Ulf Schlichtmann, Technical University of Munich, Germany

Industry Liaison

Haoxing (Mark) Ren, NVIDIA, USA

Venue Map & Information

The Westin Grand Munich

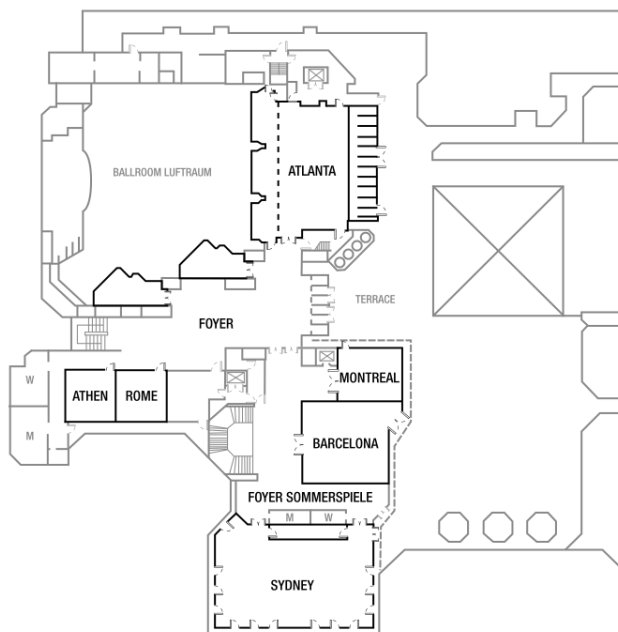
Address: Arabellastraße 6, 81925 München, Germany

Parking at the Hotel

On-Site Parking: Hourly: €3.50; Daily: €35.00

Venue Map

MEETING ROOMS





CHIP DESIGN GERMANY

THE NETWORK!

Chipdesign Germany is the leading agile network for chip design in Germany.

Chipdesign Germany

- ... strengthens microelectronics in Germany
- ... connects industry and research
- ... is the think tank for chip design
- ... attracts skilled labour
- ... increases the visibility of chip design



Contact

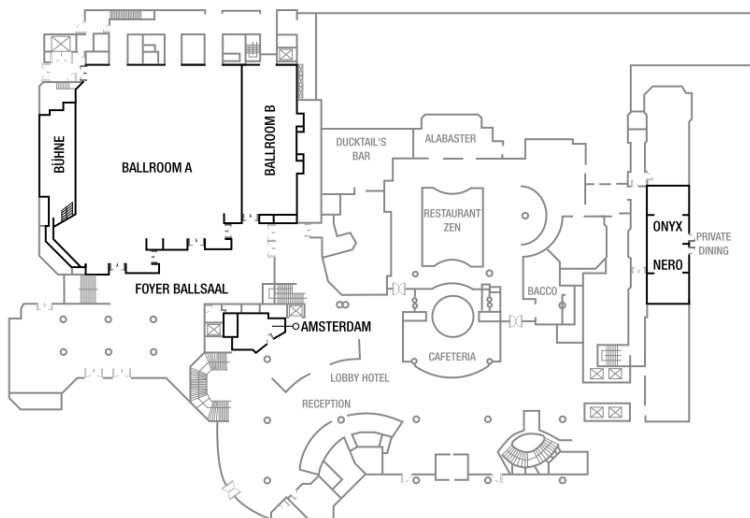
info@chipdesign-germany.de

www.chipdesign-germany.de

**Creating innovations together!
Your expertise counts!**

Venue Map & Info (cont.)

MEETING ROOMS



Welcome Message

Dear Attendees of ICCAD 2025,

Welcome to Munich, Germany—and **welcome to the 44th International Conference on Computer-Aided Design (ICCAD)**. It is our immense pleasure to host this vibrant community for the next five days and to offer an exciting program filled with opportunities for exchange, networking, and professional growth.

Jointly sponsored by ACM and IEEE, ICCAD **remains the premier venue** for presenting cutting-edge research and exploring emerging technologies in computer-aided design. From devices and circuits to architectures, systems, as well as applications—especially in the age of AI and heterogeneous computing—ICCAD brings together thought leaders, innovators, and passionate researchers from around the world.



And ICCAD continues to thrive and evolve! **For the first time in its over 40-year history, the conference is being held in Europe**—marking a major milestone and underscoring ICCAD's transformation into a truly international event. Moreover, **for the sixth consecutive year, the number of submissions has increased**. Since 2020, ICCAD has **more than doubled its submission volume**, a testament to the strength, vitality, and expanding reach of the EDA/CAD research community—and a guarantee of an outstanding technical program.

This year, ICCAD has set a new record for paper submissions, with **1508 abstract submissions across 19 technical tracks from 37 regions and countries**. A total of 1078 full papers underwent a rigorous review process, supported by 384 exceptional expert reviewers on our Technical Program Committee. Ultimately, **266 papers were accepted**, resulting in a competitive acceptance rate of 24.7%. These papers are organized into 50 sessions, **forming the largest technical program in ICCAD's history**.

In addition to the regular sessions, we are proud to feature:

- **14 special sessions** exploring the future of CAD—including microfluidics, chiplet systems, hardware-software co-design for AI, LLM foundations, neuromorphic computing, and more.
- **2 embedded tutorials** on hardware security and quantum architecture design.
- And, for the first time at ICCAD, a **panel discussion** on the limits and potential of large AI models in hardware design.

Our program is crowned by three **distinguished keynotes** from leaders in the field:

- On Monday morning, Diana Marculescu (University of Texas at Austin) will speak on “The Quest for Energy-Efficient Generative AI.”
- On Tuesday morning, Hans-Jörg Vögel (BMW Group) will present “Driven by AI – Driving Requires More Compute Than Ever.”
- On Wednesday morning, Luca Benini (Università di Bologna) will share insights on “End-to-End Open Source Platforms in the Era of Domain-Specific Design Automation.”

Welcome Message (cont.)

Finally, on Thursday, we invite you to join several **co-located workshops** covering topics such as: Automotive Chiplets, Post-Quantum Cryptography, Top Picks in Hardware & Embedded Security, Sustainable Hardware Security, System Level Interconnect Pathfinding, Quantum Computing Applications & Systems, Foundation Models and EDA. These workshops offer a space for focused discussion and deeper engagement with emerging areas of research.

In addition, throughout the week, ICCAD will also host **sessions dedicated to workforce development and student engagement**. Our integrated Student Scholar Program supports young researchers through activities such as the SIGDA CADathlon, the ACM Student Research Competition, and the SIGDA Job Fair. Thanks to generous support from our sponsors, we are proud to welcome many first-time attendees and provide travel grants to help students connect with the ICCAD community.

None of this would be possible without the tireless efforts of our volunteers—including the executive committee, technical program committee, workshop organizers, and student activity coordinators. We also want to extend **heartfelt thanks to our sponsors for their continued support** and belief in ICCAD's mission.

Finally, **the success of ICCAD hinges on your participation**. We hope you will have a great time, gain new insights, share recent accomplishments, reconnect with old friends, and forge new connections. Whether you're here to present, learn, collaborate, or simply engage with the community—welcome to ICCAD 2025! **We hope you enjoy the conference, the city of Munich, and the vibrant spirit that makes ICCAD so special.**

Warm regards,

Robert Wille
General Chair
Technical University of Munich,
Munich Quantum Software Company, and
Software Competence Center Hagenberg
GmbH

Deming Chen
Program Chair
University of Illinois Urbana-Champaign

Ismail S. K. Bustany
Vice Program Chair
AMD

Tsung-Yi Ho.png
Tutorial & Special Session Chair
The Chinese University of Hong Kong

Ron Duncan
Workshop Chair
Synopsys

Awards

William J. McCalla ICCAD Best Paper Award -

Frontend: 1239: LaZagna: An Open-Source Framework for Flexible 3D FPGA Architectural Exploration

Ismael Youssef {1}, Hang Yang {1}, Cong "Callie" Hao {2}
{1} Georgia Tech, United States; {2} Georgia Institute of Technology, United States

Backend: 582: Semidefinite Programming-Based Decoupling Capacitor Placement for Power Distribution Network Optimization

Zong-Ying Cai {1}, Wei-Han Mao {1}, Yao-Wen Chang {1}, Yang Lu {2}, Jerry Bai {2}, Bin-Chyi Tseng {2}
{1} National Taiwan University, Taiwan; {2} ASUSTeK Computer Inc., Taiwan

ICCAD 10 Year Retrospective Most Influential Paper Award

10.1145/2966986.2967011: Caffeine: Towards uniformed representation and acceleration for deep convolutional neural networks

Chen Zhang; Zhenman Fang; Peipei Zhou; Peichen Pan; Jason Cong

2025 Outstanding Service Recognition Award



Jinjun Xiong (University at Buffalo)

"for outstanding service to the EDA Community as ICCAD General Chair in 2024"

Awards (cont.)

Best Reviewer Award

BaekGyu Kim, Daegu Gyeongbuk Institute of Science and Technology, South Korea

Chen Zhang, Shanghai Jiao Tong University, China

Fan Chen, Indiana University Bloomington, USA

Tathagata Srimani, CMU, USA

Yukai Chen, IMEC, Belgium

This award recognizes the dedicated service of reviewers and the outstanding quality of reviews received by TPC members. The awardees above were selected from a pool of candidates nominated by the TPC Track co-chairs, based on the following criteria:

- **Insightfulness and Clarity:** detailed comments with strong insights, e.g., feedback is clear and specific without vagueness.
- **Thoroughness and Quality:** thorough and high quality, e.g., comments are comprehensive and address both strengths and areas for improvement.
- **Originality of Contribution:** own review instead of relying on the secondary reviewer.
- **Consistency and Fairness:** the review scores should be consistent with the review content; overall tone and recommendations align with the content of the feedback.
- **Impact on Authors:** consider whether the review would genuinely help the author improve their work (a key sign of a high-quality review).
- **Professionalism:** ensure the review is respectful, free of bias, and professionally written.

BPA Candidates (Front End)

26: GauPRE: A Pattern-based Rendering Engine for Gaussian Splatting on Edge Device

Yuzheng Lin {1}, Lizhou Wu {1}, Chixiao Chen {2}, Xiaoyang Zeng {3}, Haozhe Zhu {3}
{1} State Key Laboratory of Integrated Chips & Systems, Frontier Institute of Chip and System, Fudan University, China; {2} Fudan University, China; {3} State Key Laboratory of Integrated Chips & Systems, Fudan University, China

44: Perturbation-efficient Zeroth-order Optimization for Hardware-friendly On-device Training

Qitao Tan {1}, Sung-En Chang {2}, Rui Xia {3}, Huidong Ji {4}, Chence Yang {1}, Ci Zhang {1}, Jun Liu {2}, Zheng Zhan {2}, Zhenman Fang {5}, Zhuo Zou {4}, Yanzhi Wang {2}, Jin Lu {1}, Geng Yuan {1}
{1} University of Georgia, United States; {2} Northeastern University, United States; {3} University of Pennsylvania, United States; {4} Fudan University, China; {5} Simon Fraser University, Canada

79: ExactMap: Enhancing Delay Optimization in Parallel ASIC Technology Mapping

Zhenxuan Xie {1}, Lixin Liu {2}, Tianji Liu {1}, Evangeline Young {1}
{1} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} The Chinese University of Hong Kong, China

85: COBRA: Algorithm-Architecture Co-optimized Binary Transformer Accelerator for Edge Inference

Ye Qiao, Zhiheng Chen, Yian Wang, Yifan Zhang, Yunzhe Deng, Sitao Huang
University of California, Irvine, United States

537: Pathfinder: Constructing Cycle-accurate Taint Graphs for Analyzing Information Flow Traces

Katharina Ceesay-Seitz {1}, Flavien Solt {2}, Alexander Klukas {1}, Kaveh Razavi {1}
{1} ETH Zurich, Switzerland; {2} UC Berkeley, United States

926: Diff-DiT: Temporal Differential Accelerator for Low-bit Diffusion Transformers on FPGA

Shidi Tang {1}, Pengwei Zheng {2}, Ruiqi Chen {3}, Yuxuan Lv {2}, Bruno Silva {3}, Ming Ling {2}
{1} Southeast University, China; {2} Southeast University, China; {3} Vrije Universiteit Brussel, Belgium

1008: The Fellowship of the Leak: Power Analysis of a Masked FrodoKEM Hardware Accelerator

Martin Schmid {1}, Giuseppe Manzoni {1}, Aydin Aysu {2}, Elif Bilge Kavun {3}
{1} University of Passau, Germany; {2} North Carolina State University, United States; {3} Barkhausen Institut & TU Dresden, Germany

BPA Candidates (cont.)

1148: MIRAGE: Microarchitectural Footprints for Detecting Adversarial Attacks in One-Shot Inference

Soumi Chatterjee {1}, Debadrita Talapatra {1}, Nimish Mishra {1}, Aritra Hazra {2}, Debdeep Mukhopadhyay {3}

{1} Indian Institute of Technology Kharagpur, India; {2} Dept of CSE, IIT Kharagpur, India;

{3} Department of Computer Science and Engineering, Indian Institute of Technology Kharagpur, India

1239: LaZagna: An Open-Source Framework for Flexible 3D FPGA Architectural Exploration

Ismael Youssef {1}, Hang Yang {1}, Cong "Callie" Hao {2}

{1} Georgia Tech, United States; {2} Georgia Institute of Technology, United States

1298: 3D Acceleration for Mixture-of-Experts and Multi-Head Attention Spiking Transformers with Dynamic Head Pruning

Boxun Xu {1}, Junyoung Hwang {2}, Pruek Vanna-iampikul {3}, Yuxuan Yin {1}, Sung Kyu Lim {4}, Peng Li {1}

{1} University of California, Santa Barbara, United States; {2} Georgia Institute of Technology, United States; {3} Burapha University, Thailand; {4} Georgia Tech, United States

1318: e-boost: Boosted E-Graph Extraction with Adaptive Heuristics and Exact Solving

Jiaqi Yin {1}, Zhan Song {1}, Chen Chen {1}, Yaohui Cai {2}, Zhiru Zhang {2}, Cunxi Yu {1}

{1} University of Maryland, College Park, United States; {2} Cornell University, United States

BPA Candidates (Back End)

95: GTA: GPU-Accelerated Track Assignment with Lightweight Lookup Table for Conflict Detection

Chunyuan Zhao, Jiarui Wang, Xun Jiang, Jincheng Lou, Yibo Lin
Peking University, China

265: CLASS: A Controller-Centric Layout Synthesizer for Dynamic Quantum Circuits

Yu Chen {1}, Yilun Zhao {2}, Bing Li {3}, He Li {4}, Mengdi Wang {1}, Yinhe Han {5}, Ying Wang {6}

{1} Institute of Computing Technology, Chinese Academy of Sciences, China; {2} Institute of Computing Technology, CAS, China; {3} Capital Normal University, China; {4} Southeast University, China; {5} Institute of Computing Technology, Chinese Academy of Sciences, China; {6} State Key Laboratory of Computer Architecture, Institute of Computing Technology, Chinese Academy of Sciences, China

554: LMLitho: A Large Vision Model-Driven Lithography Simulation Framework

Zhen Wang {1}, Hongquan He {1}, Tao Wu {1}, Xuming He {2}, Qi Sun {3}, Cheng Zhuo {3}, Bei Yu {4}, Jingyi Yu {1}, Hao Geng {1}

{1} ShanghaiTech University, China; {2} ShanghaiTech University, China; {3} Zhejiang University, China; {4} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China

561: FAB: Fast and Demand-Aware Bandwidth Allocation Method for Wavelength-Routed Optical Networks-on-Chip

Liaoyuan Cheng, Mengchu Li, Zhidan Zheng, Tsun-Ming Tseng, Ulf Schlichtmann
Technical University of Munich, Germany

582: Semidefinite Programming-Based Decoupling Capacitor Placement for Power Distribution Network Optimization

Zong-Ying Cai {1}, Wei-Han Mao {1}, Yao-Wen Chang {1}, Yang Lu {2}, Jerry Bai {2}, Bin-Chyi Tseng {2}

{1} National Taiwan University, Taiwan; {2} ASUSTeK Computer Inc., Taiwan

821: TickTockStack: In-Datapath Current Imbalance Elimination Using Clocked Differential Logic in a Voltage Stacked Vector Processor

Michal Gorywoda, Wanyoung Jung
KAIST, Republic of Korea

860: Leveraging GPU for Better Detailed Placement Quality

Chen-Han Lu {1}, Wen-Hao Liu {2}, Haoxing Ren {3}, Ting-Chi Wang {1}

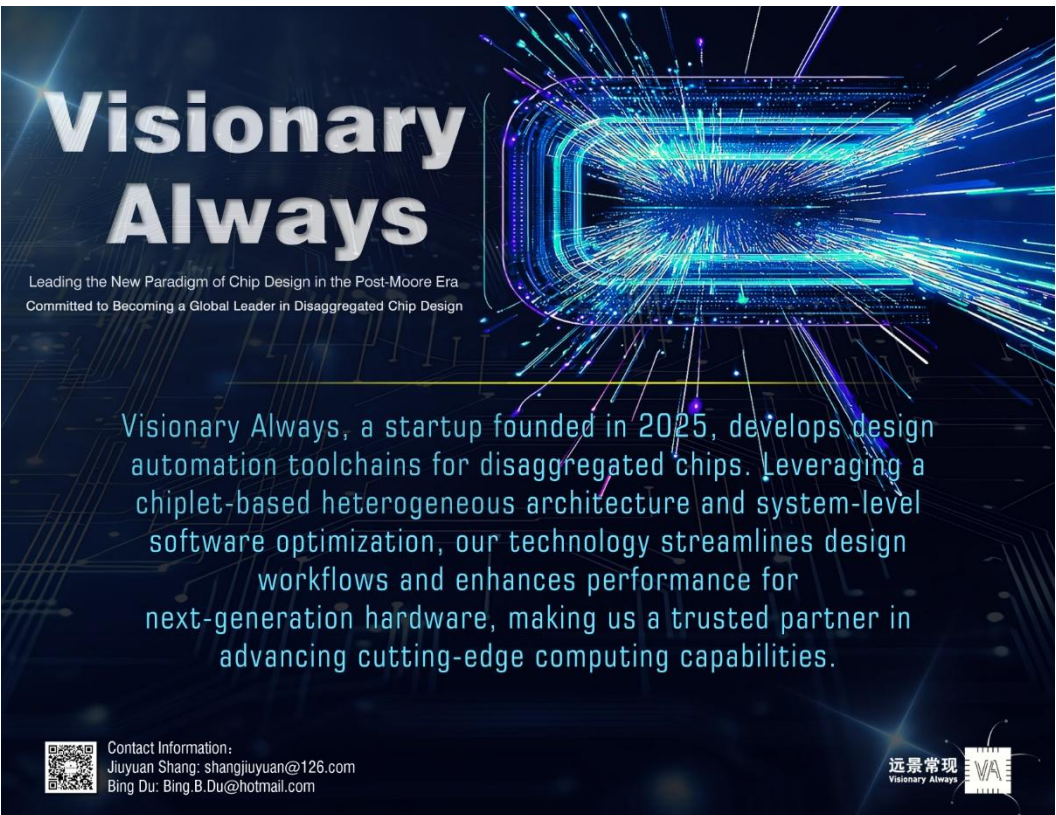
{1} National Tsing Hua University, Taiwan; {2} Nvidia, Taiwan; {3} NVIDIA Corporation, United States

BPA Candidates (cont.)

1231: dyGRASS: Dynamic Spectral Graph Sparsification via Localized Random Walks on GPUs

Yihang Yuan {1}, Ali Aghdaei {2}, Zhuo Feng {1}

{1} Stevens Institute of Technology, United States; {2} University of California San Diego, United States

The advertisement features a dark blue background with a glowing, futuristic circuit pattern. On the right side, there is a large, glowing blue and purple rectangular frame that resembles a tunnel or a portal, with many bright light streaks radiating from its center. The text "Visionary Always" is prominently displayed in a large, white, sans-serif font on the left. Below it, in a smaller white font, is the tagline "Leading the New Paradigm of Chip Design in the Post-Moore Era. Committed to Becoming a Global Leader in Disaggregated Chip Design." Further down, a paragraph in a light blue font describes the company's mission and technology. At the bottom left, there is a QR code and contact information. At the bottom right, there is a logo with the Chinese characters "远景常现" and the English text "Visionary Always" next to a stylized "VA" monogram.

Visionary Always

Leading the New Paradigm of Chip Design in the Post-Moore Era.
Committed to Becoming a Global Leader in Disaggregated Chip Design

Visionary Always, a startup founded in 2025, develops design automation toolchains for disaggregated chips. Leveraging a chiplet-based heterogeneous architecture and system-level software optimization, our technology streamlines design workflows and enhances performance for next-generation hardware, making us a trusted partner in advancing cutting-edge computing capabilities.

Contact Information:
Jiuyuan Shang: shangjiuyuan@126.com
Bing Du: Bing.B.Du@hotmail.com

远景常现
Visionary Always

cādence®

Enabling the Design of the AI World

Scan the QR to learn more!



Visit the Cadence exhibit table at ICCAD

Plenary Speakers



The Quest for Energy Efficient Generative AI

Monday | October 27, 2025 | 8:30 - 9:30

Room: Ballroom A+B

Diana Marculescu, University of Texas at Austin

Abstract: Artificial Intelligence (AI) applications have entered and impacted our lives unlike any other technology advance from the recent past. While current

AI tools are approaching Artificial General Intelligence (AGI) capabilities, large gaps persist due to diminished benefits despite deploying enormous computing and ever-increasing energy demands. This quandary requires a holistic approach in determining energy efficient solutions for achieving AGI. While the holy grail for judging the quality of an AI model has largely been serving accuracy, and only recently its resource usage, neither of these metrics translate directly to energy efficiency, latency, or mobile device battery lifetime. More recently, generative AI relying on transformer, diffusion, or state-space models have revolutionized the way we approach reasoning and learning tasks across all types of modalities, from image to video and language, bringing increased performance at the expense of significant hardware costs.



Driven by AI – Driving Requires More Compute Than Ever

Tuesday | October 28, 2025 | 8:00 - 9:00

Room: Ballroom A+B

Hans-Jörg Vögel, BMW Group

Abstract: Higher levels of automation in road transport mobility are a true moonshot. Very few OEMs have so far

successfully obtained road approval for their Level 3 automated driving system, BMW being one of them. AI already now is a major technological driving force behind automated driving and will continue to be. We will highlight current frontiers of AI development in safety relevant systems and explore, how expected evolution of the technology will influence the road ahead for automated driving functions. Growth in sensor data volume and machine learning compute demand is fuelling a race to ever more potent embedded systems. Just brute-forcing hardware will not do. Innovative semiconductor solutions such as chiplets come to the rescue. As does intelligent hardware-software-co-design.

Plenary Speakers (cont.)



End-to-end Open Source Platforms in the era of Domain-Specific Design Automation

Wednesday | October 29, 2025 | 8:00 - 9:00

Room: Ballroom A+B

Luca Benini, Università di Bologna

Abstract: Modern computing workloads, from generative AI, to digital twins, to fully homomorphic encryption, require sustained energy efficiency improvements, which cannot be matched simply by technology evolution. To tackle the challenge, we need to aggressively optimize computing platforms leveraging specialization across all levels of the design hierarchy, pushing into domain-specific design automation tools and methodologies. In this talk, I will give concrete examples of deep domain specialization, emphasizing the strategic importance of an end-to-end (models, software, instruction set architecture, digital IPs, EDA tools) open-platform approach to achieve the ultimate efficiency, while promoting a healthy innovation ecosystem.

Tutorial Sessions

Silence of the Chips: Advanced Techniques in Hardware Vulnerability Detection

Monday | October 27, 2025 | 10:00 - 11:30

The ever-increasing complexity of microprocessors has resulted in several potent security threats in recent years. These vulnerabilities in the hardware arising from unchecked performance optimizations have been exploited through various ways, such as micro-architectural attacks, fault injections, memory corruption, and other forms of information leakage. Post tape-out, hardware vulnerabilities are typically mitigated using software updates or hardware recall, which result in unacceptably high performance or economic overheads, respectively. Thus, there is a pressing need to uncover these vulnerabilities during the hardware design phase. Integrating such approaches can improve the overall security, reliability, and economic viability of microprocessors. In this tutorial, we introduce hardware vulnerabilities and state-of-the-art techniques to uncover these vulnerabilities at design time, leveraging hardware fuzzing, AI, and formal verification. Tutorial participants will gain an understanding of the fundamentals of hardware vulnerabilities, their origins, and detection approaches. We will present some of the potent Common Weakness Enumerations (CWEs) we have exposed in popular microprocessors. With our assistance, participants will get a real-world demonstration of the hardware fuzzing techniques used to detect these vulnerabilities and pinpoint their location in hardware design. We will explore the recent advances in hardware fuzzing using AI techniques and formal verification. We will use concrete, hands-on examples to quantitatively analyze the potential of these techniques for hardware security and the open challenges in the domain.



Ahmad-Reza
Sadeghi
TU Darmstadt



JV Rajendran
Texas A&M
University



Nikhilesh Singh
TU Darmstadt



Lichao Wu
TU Darmstadt



Huimin Li
TU Darmstadt



Chen Chen
Texas A&M
University



Mohamadreza
Rostami
TU Darmstadt

Tutorial Sessions (cont.)

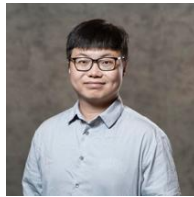
Quantum Machine Learning Foundations for Quantum Architecture Design and Future Quantum EDA

Monday | October 27, 2025 | 13:00 - 14:30

This tutorial explores the intersection of quantum machine learning (QML) and the design of quantum architectures, highlighting foundational techniques that pave the way for future quantum-enabled design automation (Quantum EDA). Rather than optimizing classical semiconductor flows, we focus on how QML, quantum reinforcement learning (QRL), and quantum architecture search (QAS) can be applied to discover, optimize, and innovate quantum circuits and variational architectures. Through interactive, hands-on sessions using open-source quantum simulators such as Qiskit and PennyLane, participants will engage with key methodologies for quantum model design and architectural optimization. This session is aimed at future explorers — those who seek to rethink design automation in the quantum era — and no prior knowledge of traditional EDA is required.



Samuel Yen-Chi Chen
Wells Fargo



Zhiding Liang
Rensselaer Polytechnic
Institute



Kuan-Chen Chen
Imperial College London

Workshops

Top Picks in Hardware and Embedded Security

Thursday | October 30, 2025 | 8:00 - 16:30

Top Picks Workshop creates a venue to showcase the best and high impact recently published works in the area of hardware and embedded security. These works will be selected from conference papers that have appeared in leading hardware security conferences including but not limited to DAC, ICCAD, DATE, ASPDAC, HOST, Asian HOST, GLSVLSI, VLSI Design, CHES, ETS, VTS, ITC, S&P, Usenix Security, CCS, NDSS, ISCA, MICRO, ASPLOS, HPCA, HASP, ACSAC, Euro S&P, and Asia CCS. The 8th Top Picks workshop will be collocated with ICCAD 2025. The authors of a short list of papers picked from the submissions are required to present their work at the workshop on October 30, 2025 and are invited to attend ICCAD's networking in-person on October 30, 2025. The presentation including a discussion (Q&A) session, is mandatory.



Ahmad-Reza Sadeghi
Technical University of Darmstadt
Darmstadt, Hesse, Germany



Xiaolin Xu
Northeastern University, USA

Workshops (cont.)

Workshop on Post-Quantum Cryptography Resilience, Verification, and Secure Design Automation (WPQC)

Thursday | October 30, 2025 | 8:00 - 16:30

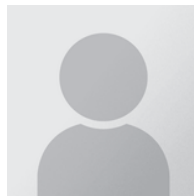
The rapid transition toward Post-Quantum Cryptography (PQC) has prompted a global effort to redefine digital security in anticipation of quantum-capable adversaries. While PQC algorithms are mathematically robust, their secure and efficient integration into hardware systems introduces new vulnerabilities and engineering challenges - especially in areas such as side-channel resistance, hardware/software co-design, and formal verification. The Workshop on Post-Quantum Cryptography and Secure Hardware (WPQC) aims to address these challenges by serving as a dedicated forum for researchers and practitioners in hardware security, EDA/CAD, computer architecture, and cryptography. The workshop provides a platform to explore the intersection of PQC and hardware/system design, highlighting recent advancements in secure accelerator development, design automation, system integration, and emerging threat models. By bringing together experts from academia, industry, and standardization bodies, WPQC seeks to accelerate collaborative research efforts and foster a shared vision for building resilient, verifiable, and efficient post-quantum systems.



Shivam Bhasin
NTU (Singapore)



Trevor E. Carlson
NUS (Singapore)



Mona Hashemi
NUS (Singapore)



Reza Azarderakhsh
FL Atlantic U, and
PQSecure (USA)



Mojtaba Bisheh-
niasar
Microsoft (USA)

Workshops (cont.)

2nd Quantum Computing Applications and Systems (QCAS)

Thursday | October 30, 2025 | 8:00 - 16:30

Our proposed workshop aims to address emerging challenges and explore innovative solutions in the field of quantum technologies, particularly focusing on quantum computing applications for real-world problems. Quantum technologies are becoming crucial in a variety of scientific domains including chemistry, finance, power systems, etc. Our workshop will bring together researchers, practitioners, and industry experts to exchange ideas, share applications, and discuss the latest advancements in quantum algorithms, error correction, control techniques, and their applications across diverse fields. This event will serve as a platform to showcase cutting-edge research, foster collaboration, and drive innovation in the intersection of CAD, optimization, machine learning, and quantum computing.



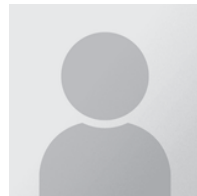
Robert Wille
TUM, MQSC, SCCH



Martin Schulz
Technical
University of
Munich, Germany



Samuel Yen-Chi
Chen
Wells Fargo



Yosuke Ueno
Riken

Workshops (cont.)

System-Level Interconnect Pathfinding (SLIP)

Thursday | October 30, 2025 | 8:00 - 16:30

SLIP, co-located with ICCAD, brings together researchers and practitioners who have a shared interest in the challenges and futures of system-level interconnect, coming from wide-ranging backgrounds that span system, application, design and technology.

The technical goal of the workshop is to

- identify fundamental problems, and,
- foster new pathfinding of design, analysis, and optimization of system-level interconnects with emphasis on system-level interconnect modeling and pathfinding, DTCO-enhanced interconnect fabrics, memory and processor communication links, novel dataflow mapping for machine learning, 2.5/3D architectures, and new fabrics for the beyond-Moore era.

Original submissions in the form of regular technical papers, invited sessions (tutorials, panels, special-topic sessions), workshop discussion topics, and posters are welcome. Program content is accepted based on novelty and contributions to the advancement of the field. Authors will keep the copyright of their work and there will be no published proceedings.

Workshops (cont.)

Foundation Models and EDA

Thursday | October 30, 2025 | 8:00 - 11:30

The rapid advancement of foundation models has brought powerful new capabilities to Electronic Design Automation (EDA). Unlike traditional task-specific Artificial Intelligence (AI) approaches, foundation models leverage self-supervised learning and large-scale data pre-training to acquire strong generalization abilities. These models can then be efficiently fine-tuned on EDA-specific datasets, enabling a wide range of downstream applications such as hardware code generation, debug and optimization, EDA agents, circuit representation learning and understanding, etc.

This workshop aims to foster collaboration among researchers, engineers, and industry experts to explore the evolving intersection of foundation models and EDA methodologies. By identifying key challenges and exchanging ideas, we seek to inspire comparative analysis between the unique demands of EDA and those of other application domains, and to promote novel solutions that leverage the strengths of foundation models for more accurate, efficient, and scalable design automation. Through talks, paper presentations, and interactive discussions, the workshop will showcase recent progress and explore new ways to apply foundation models in EDA—helping to drive innovation and advance the future of intelligent circuit design.



Deming Chen
University of
Illinois



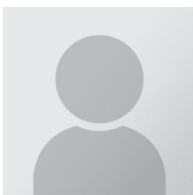
Igor Markov
Synopsys



Guohao Dai
Shanghai Jiao Tong
University



Siddharth Garg
NY University



Zhengyuan Shi
Chinese University
of Hong Kong

Workshops (cont.)

Automotive Chiplet Platform: Solutions Enabled by Industry and Academic Collaboration

Thursday | October 30, 2025 | 13:00 - 16:30

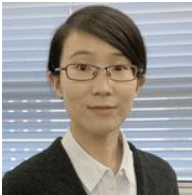
Chiplet-based platform is a pragmatic solution to meet the automotive industry's growing demands for performance, longevity, and cost-efficiency. Such platforms enable modular, scalable integration of diverse processing units (like CPUs, GPUs, and AI accelerators) across different process nodes, improving design flexibility, yield, and reliability, which is critical in safety-driven environments. They also support thermal and power optimization, simplify long-term maintenance and supply chain challenges, and align well with emerging zonal E/E architectures. The purpose of the workshop is to outline the challenges and novel solutions to enable chiplet-based computing platform for automotive domain.

Panel Session

Revolution or Hype? Seeking the Limits of Large Models in Hardware Design

Tuesday | October 28, 2025 | 9:30 - 11:00

Recent breakthroughs in Large Language Models (LLMs) and Large Circuit Models (LCMs) have sparked excitement across the electronic design automation (EDA) community, promising significant advances in circuit design and optimization. Yet, skepticism persists: Are these large AI models truly transformative, or are we experiencing a temporary wave of inflated expectations? This panel brings together leading experts from academia and industry to critically examine the practical capabilities, limitations, and future prospects of large AI models in hardware design. Panelists will debate their reliability, scalability, interpretability, and whether these models meaningfully outperform and/or complement traditional EDA methods—offering attendees fresh insights into one of today's most contentious and impactful technology trends.



Moderator
Grace Li Zhang
TU Darmstadt



Xi Wang
Southeast University



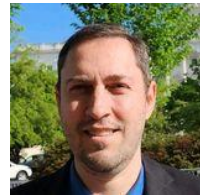
Qiang Xu
National Technology
Innovation Center for
EDA, China



Rolf Drechsler
University of Bremen



Leon Stok
IBM



Igor L. Markov
Synopsys

CEDA Luncheon

From Past to Future: 20 Years of EDA Wisdom and What's Next

Monday | October 27, 2025 | 11:30 - 13:00

Over the past two decades, the EDA community and industry have made significant contributions to enabling productive chip designs and driving the overall progress and growth of semiconductor technology. Looking ahead, EDA faces ongoing challenges and opportunities in a rapidly evolving landscape, largely fueled by the phenomenal advances in ML/AI technologies. How can we seize these opportunities while addressing the ever-growing complexity of chip design and the relentless demands of the market? In this panel, four distinguished experts from academia and industry around the world will share insights on lessons learned from the past and their visions for the future.

Panelists



Subhasish Mitra
Stanford
University, USA



Shaojun Wei
Tsinghua
University, China



Sani Nassif
Radyalis, USA



Ulf Schlichtmann,
Technical
University of
Munich, Germany

Moderators



L. Miguel Silveira
IEEE CEDA
President



Robert Wille
TUM, MQSC, SCCH

Workforce Development

Student Research Competition

Posters: Monday | October 27, 2025 | 19:00 - 20:30

Oral Presentations: Tuesday | October 28, 2025 | 12:30 - 14:00

The ACM Student Research Competition (SRC) provides undergraduate and graduate students who are ACM members with a unique opportunity to present their work, exchange ideas, and gain recognition from the international computing community. Organized by ACM SIGDA and hosted at ICCAD 2025, the competition features poster and presentation sessions where students showcase their research in design automation. Participants benefit from feedback by experts, networking with academic and industry leaders, and the chance to win prizes. Separate categories for undergraduate and graduate students ensure fair evaluation, with submissions judged on novelty, technical merit, and contributions to the field.

Job Fair

Tuesday | October 28, 2025 | 18:00 - 19:00

ACM SIGDA and IEEE CEDA will host the 4th annual EDA Job Fair at ICCAD 2025, providing students and professionals the opportunity to connect with leading companies and academic institutions. Attendees can learn about current openings, network directly with recruiters, and explore career opportunities in electronic design automation. Participants are encouraged to submit their CVs in advance, and ICCAD registration (minimum one-day) is required.

CADathlon@ICCAD 2025

Sunday | October 26, 2025 | 8:00 - 17:00

The CADathlon is a challenging, all-day programming competition focusing on practical problems at the forefront of Computer-Aided Design and Electronic Design Automation in particular. The contest emphasizes the knowledge of algorithmic techniques for CAD applications, problem-solving and programming skills, and teamwork.

As the “Olympic games of EDA,” the contest brings together the best and the brightest of the next generation of CAD professionals. It gives academia and the industry a unique perspective on challenging problems and rising stars, and it also helps attract top graduate students to the EDA field.

Social Functions

Welcome Reception

Monday | October 27, 2025 | 19:00 - 21:30 | The Westin Grand Munich

All conference registrants are invited to join us for the Welcome Reception on Monday, October 27, at The Westin Grand Munich. Start the conference on a high note with an evening of networking, hors d'oeuvres, and refreshments as you connect with fellow attendees in a relaxed and friendly setting. We look forward to welcoming you in Munich!

ACM SIGDA Gala Dinner

Wednesday | October 29, 2025 | 18:30 - 20:00 | Hofbräuhaus München Festsaal

Join us for an unforgettable evening at the historic Hofbräuhaus München Festsaal, featuring a nine-meter-high barrel vault and classic Bavarian ambiance. Enjoy traditional food, drinks, and lively conversation in one of Munich's most iconic venues. Access is included with your conference registration—no additional ticket required. Attendees are responsible for their own transportation to and from the venue.

Address: Platzl 9, 80331 München



Sponsors & Exhibitors

Platinum Sponsor

SYNOPSYS®

ACADEMIC & RESEARCH
ALLIANCES

Gold Sponsors

cādence®



Silver Sponsors

AMD



Sony AI



Open Innovation Platform®



Sponsors & Exhibitors

Conference Sponsors



Sunday, October 26, 2025

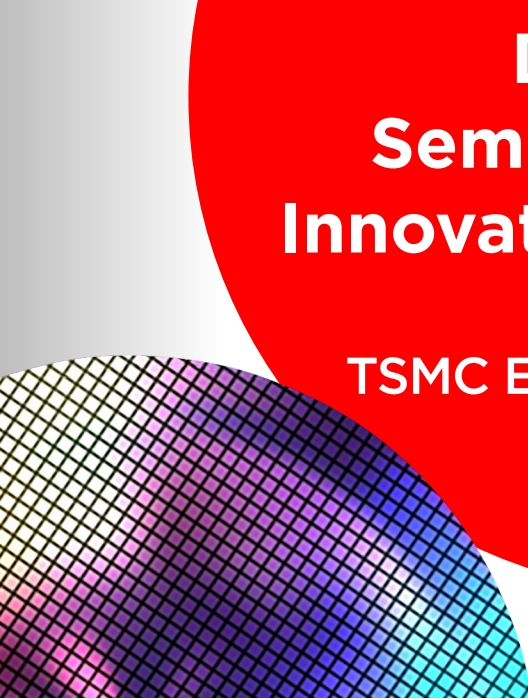
7:30 - 8:00	Registration Room: Foyer Ballroom
8:00 - 9:00	CADathlon Breakfast Room: Atlanta
9:00 - 12:00	CADathlon Room: Atlanta
12:00 - 13:00	CADathlon Lunch Room: Atlanta
13:00 - 15:00	CADathlon Room: Atlanta
15:00 - 15:30	CADathlon Coffee Break Room: Atlanta
15:30 - 17:00	CADathlon Room: Atlanta

Monday, October 27, 2025						
Time	Sydney	Barcelona	Atlanta	Ballroom A	Ballroom B	Calgary
7:00 - 8:00	Registration Room: Foyer Ballroom					
8:00 - 8:30	Opening Ceremony Room: Ballroom A					
8:30 - 9:30	Keynote: Diana Marculescu Room: Ballroom A					
9:30 -10:00	Coffee Break Exhibits Room: Foyer Ballroom					
10:00 - 11:30	Tutorial: Silence of the Chips: Advanced Techniques in Hardware Vulnerability Detection	Technical Session: Speeding and Learning for Achieving 2D and 3D Design Closure	Special Session: Agile and Open Hardware Design and Verification	Technical Session: Hardware Assurance in the Adversarial Era: Camouflage, Detection, and Lightweight Defense	Technical Session: Memory and Storage Design Optimization Techniques	Technical Session: Journey from Efficient Layout to Successful Wafer
11:30 - 13:00		CEDA Luncheon Panel Room: Ballroom A				
13:00 - 14:30	Tutorial: Quantum Machine Learning Foundations for Quantum Architecture Design and Future Quantum EDA	Technical Session: Processing-in-Memory (PIM/ CIM) Powered AI Acceleration	Special Session: Hardware-Software Co-Design for highly Optimized, Customized, and Reliable AI Systems	Technical Session: Logic Synthesis Reimagined: Mapping, Rewriting, and Structural Optimization	Technical Session: Hardware Reliability and Testing	Technical Session: Silicon Whispers: Side Channels, Rowhammer, and Microarchitectural Leaks
14:30 - 15:00		Coffee Break Exhibits Room: Foyer Ballroom				
15:00 - 16:30	Special Session: EDA from GPU Acceleration to GPU Exploration	Technical Session: Edge Intelligence: Co-Design, Resilience, and Efficiency	Special Session: Security Challenges and Opportunities in In-Sensor Computing Systems	Technical Session: Power and Performance Conscious Hardware Design	Technical Session: Advanced Modeling and Optimization for Power and Thermal Integrity	Technical Session: In memory computing and computing in memory
16:30 - 16:45	Transition					
16:45 - 17:45	Technical Session: Edge-Intelligent: Co-Design and Compression for Efficient AI at the Edge	Technical Session: Advancing towards the Optimum Cell Generation	Technical Session: Learning-Based Hardware Optimization and Scheduling	Technical Session: Timing Optimization and Optical Routing	Technical Session: Optimization and Simulation From Circuits to Systems	Technical Session: Emerging Architectures and Automation for Neuromorphic, In-Memory, and Microfluidic Computing
18:00 - 19:00	Synopsis Sponsor Session Room: Ballroom A					
19:00 - 20:30	Welcome Reception & SRC Poster Session Room: Foyer Ballroom					

Tuesday, October 28, 2025						
Time	Sydney	Barcelona	Atlanta	Ballroom A	Ballroom B	Calgary
7:30 - 8:00	Registration Room: Foyer Ballroom					
8:00 - 9:00	Keynote: Hans-Jörg Vogel Room: Ballroom A					
9:00 - 9:30	Coffee Break Exhibits Room: Foyer Ballroom					
9:30 - 11:00	Panel: Revolution or Hype? Seeking the Limits of Large Models in Hardware Design	Technical Session: AI Analog Circuits Design with Learning and Reasoning	Special Session: Circuit and Architecture Design with Emerging Computing Paradigms	Technical Session: Perfecting Routing by AI Means	Technical Session: Accelerating Large Language Models (LLMs)	Technical Session: Breaking the Thermal Wall: Emerging Thermal Simulation and Design for 2.5D/3D Systems
		Technical Session: How to Become a Better Physical Designer: Novel Tweaks on Mapping, Partitioning and Placement	Special Session: Generative AI for At-Scale and Unconventional Analog/RFIC Design and Automation	Technical Session: Addressing noise and optimization of quantum and photonic systems	Technical Session: Reasoning and Debugging in Hardware Design with LLMs	Technical Session: Brains and Bits: FPGA Acceleration from AI to SAT Solvers
11:00 - 12:30	Lunch Room: Ballroom A					
12:30 - 14:00	Student Research Competition (Oral Presentations)					
14:00 - 14:30	Coffee Break Exhibits Room: Foyer Ballroom					
14:30 - 16:00	Special Session: Fueling the Future: Achieving Reliable and Generalizable Data Foundations for LLMs in EDA	Technical Session: FPGA and Reconfigurable Accelerators for AI Models	Special Session: Reimagining Scientific Computing with Neuromorphic and Analog Systems	Technical Session: System-level and Architecture Design Optimisation	Technical Session: Next-Gen Acceleration: Sparsity, Vectors, and Reconfigurable Intelligence	Technical Session: Efficient Verification and Debugging
16:00 - 16:15	Transition					
16:15 - 17:15	Special Session: Fueling the Future: Achieving Reliable and Generalizable Data Foundations for LLMs in EDA	Technical Session: FPGA and Reconfigurable Accelerators for AI Models	Special Session: Reimagining Scientific Computing with Neuromorphic and Analog Systems	Technical Session: System-level and Architecture Design Optimisation	Technical Session: Next-Gen Acceleration: Sparsity, Vectors, and Reconfigurable Intelligence	Technical Session: Efficient Verification and Debugging
18:00 - 19:00	Job Fair Room: Foyer Ballroom					

Wednesday, October 29, 2025				
Time	Sydney	Barcelona	Atlanta	Ballroom A Ballroom B
7:30 - 8:00	Registration Room: Foyer Ballroom			
8:00 - 9:00	Keynote: Luca Benini Room: Ballroom A			
9:00 - 9:30	Coffee Break Exhibits Room: Foyer Ballroom			
9:30 - 11:00	Special Session: Cross-Layer Design for Health-Centric AI Systems: Models, Platforms, and Emerging Hardware	Technical Session: Multimodal and Generative Intelligence for EDA	Special Session: System-Level Design for Chiplets: Architecture Exploration and Resource Management	Technical Session: Flow Forward: Frameworks and Optimizations for Reconfigurable Accelerators
11:00 - 12:30	Lunch Room: Ballroom A			
12:30 - 14:00	Special Session: 2025 CAD Contest at ICCAD	Technical Session: Hardware Support for Security: Confidential Computing with Accelerators	Technical Session: Scalable Simulation and Verification Frameworks	Technical Session: From Prediction to Closure: Intelligent Techniques for Timing Estimation and Optimization
14:00 - 14:30	Coffee Break Exhibits Room: Foyer Ballroom			
14:30 - 16:00	Special Session: Privacy-Preserving Deep Learning in the LLM Era: Algorithm, Protocol, and Hardware	Technical Session: Hardware Support for Security: Confidential Computing with Accelerators	Technical Session: Scalable Simulation and Verification Frameworks	Technical Session: From Prediction to Closure: Intelligent Techniques for Timing Estimation and Optimization
16:00 - 16:15	Transition			
16:15 - 17:15	Special Session: Bridging the Gap: Design Automation Meets Microfluidics	Technical Session: Memory-Centric Systems: Design and Optimizations	Technical Session: Multi-Physics Modelling and Advanced Integration	Technical Session: Hardware for Neural Rendering and Geometric Models
18:30 - 20:00	ACM SIGDA Dinner Location: Hofbräuhaus München Festsaal			

Thursday, October 30, 2025						
Time	Sydney	Barcelona	Atlanta	Athen	Rome	Montreal
7:30 - 8:00	Registration Room: Upstairs Foyer					
8:00 - 9:30	Workshop on Post-Quantum Cryptography Resilience, Verification, and Secure Design Automation (WPQC)	Top Picks in Hardware and Embedded Security		2nd Quantum Computing Applications and Systems (QCAS) Workshop	System-Level Interconnect Pathfinding (SLIP)	Foundation Models and EDA
9:30 - 10:00	Coffee Break Room: Upstairs Foyer					
10:00 - 11:30	Workshop on Post-Quantum Cryptography Resilience, Verification, and Secure Design Automation (WPQC)	Top Picks in Hardware and Embedded Security		2nd Quantum Computing Applications and Systems (QCAS) Workshop	System-Level Interconnect Pathfinding (SLIP)	Foundation Models and EDA
11:30 - 13:00	Standing Lunch Room: Foyer Area of Sydney and Atlanta					
13:00 - 14:30	Workshop on Post-Quantum Cryptography Resilience, Verification, and Secure Design Automation (WPQC)	Top Picks in Hardware and Embedded Security		2nd Quantum Computing Applications and Systems (QCAS) Workshop	System-Level Interconnect Pathfinding (SLIP)	Automotive Chiplet Platform: Solutions Enabled by Industry and Academic Collaboration
14:30 - 15:00	Coffee Break Room: Upstairs Foyer					
15:00 - 16:30	Workshop on Post-Quantum Cryptography Resilience, Verification, and Secure Design Automation (WPQC)	Top Picks in Hardware and Embedded Security		2nd Quantum Computing Applications and Systems (QCAS) Workshop	System-Level Interconnect Pathfinding (SLIP)	Automotive Chiplet Platform: Solutions Enabled by Industry and Academic Collaboration



Driving Semiconductor Innovation in the EU

Join our
TSMC EU Design Center
now

tsmc

TSMC's European Design Center (EUDC): Driving Semiconductor Innovation in the EU

Taiwan Semiconductor Manufacturing Company (TSMC) has announced the establishment of its European Design Center (EUDC) in Munich, Germany, as part of its global network of design hubs, which includes facilities in Taiwan, the United States, and Japan. The EUDC is set to open in Q3 this year and will focus on supporting TSMC's global customers in designing high-density, high-performance, and energy-efficient chips. The center's key emphasis will be on sectors like automotive, industrial applications, artificial intelligence (AI), telecommunications, and the Internet of Things (IoT).

Munich was strategically chosen due to its proximity to a strong talent pool in the automotive industry and regional customers, including tier-1 suppliers and original equipment manufacturers (OEMs). The city provides an optimal environment for innovation and deep collaboration, aligning seamlessly with TSMC's broader strategy of enabling next-generation technology development while supporting customer needs. The EUDC also supports the EU's semiconductor strategy as outlined in the European Chips Act, contributing to building a robust semiconductor ecosystem and bolstering supply chain resilience.

The EUDC will offer significant benefits to TSMC's customers, including close collaboration to optimize chip designs for better performance, energy efficiency, and manufacturing processes. It will accelerate time-to-production and enable a faster yield learning curve, while providing system-level expertise to expand Design Technology Co-Optimization (DTCO) opportunities. Industries like automotive and emerging technologies, such as AI and IoT, will benefit from advancements in applications like Non-Volatile Memory (NVM), Resistive Random Access Memory (RRAM), and Magnetoresistive Random Access Memory (MRAM). These alternatives are increasingly important due to the rising costs of traditional eFlash technology beyond N28 nodes.

The EUDC's long-term goals include cultivating a talent pool specializing in automotive and non-volatile memory development, driving system-level optimization, and scaling innovation across various sectors. Initially, TSMC plans to hire up to 30 experienced engineers over the next few years to support these efforts. This initiative reflects TSMC's commitment to driving technological excellence globally and deepening engagement with the European market. The EUDC complements TSMC's manufacturing expansion in the EU, including the recently announced facility in Dresden, which aims to bring manufacturing capabilities closer to European customers. Through such investments, TSMC reinforces its leadership in semiconductor development while fostering innovation and collaboration on a global scale.

Technical Program – Oct. 26

7:30 – 8:00

CADathlon ONLY Registration

Room: Foyer Ballroom

8:00 – 9:00

CADathlon Breakfast

Room: Atlanta

9:00 – 12:00

CADathlon

Room: Atlanta

12:00 – 13:00

Lunch

Room: Atlanta

13:00 – 15:00

CADathlon

Room: Atlanta

15:00 – 15:30

Coffee Break

Room: Atlanta

15:30 – 17:00

CADathlon

Room: Atlanta

The CADathlon is a challenging, all-day programming competition focusing on practical problems at the forefront of Computer-Aided Design and Electronic Design Automation in particular. The contest emphasizes the knowledge of algorithmic techniques for CAD applications, problem-solving and programming skills, and teamwork.

As the “Olympic games of EDA,” the contest brings together the best and the brightest of the next generation of CAD professionals. It gives academia and the industry a unique perspective on challenging problems and rising stars, and it also helps attract top graduate students to the EDA field.

Technical Program – Oct. 27

7:00 – 8:00

Registration

Room: Foyer Ballroom

8:00 – 8:30

Opening Ceremony

Room: Ballroom A

8:30 – 9:30

Keynote: The Quest for Energy Efficient Generative AI

Diana Marculescu, University of Texas at Austin

Room: Ballroom A

Session Chair(s): Robert Wille, TUM, MQSC, SCCH

9:30 – 10:00

Coffee Break and Exhibits

Room: Foyer Ballroom

10:00 – 11:30

Tutorial: Silence of the Chips: Advanced Techniques in Hardware Vulnerability Detection

Room: Sydney

The ever-increasing complexity of microprocessors has resulted in several potent security threats in recent years. These vulnerabilities in the hardware arising from unchecked performance optimizations have been exploited through various ways, such as micro-architectural attacks, fault injections, memory corruption, and other forms of information leakage. Post tape-out, hardware vulnerabilities are typically mitigated using software updates or hardware recall, which result in unacceptably high performance or economic overheads, respectively. Thus, there is a pressing need to uncover these vulnerabilities during the hardware design phase. Integrating such approaches can improve the overall security, reliability, and economic viability of microprocessors. In this tutorial, we introduce hardware vulnerabilities and state-of-the-art techniques to uncover these vulnerabilities at design time, leveraging hardware fuzzing, AI, and formal verification. Tutorial participants will gain an understanding of the fundamentals of hardware vulnerabilities, their origins, and detection approaches. We will present some of the potent Common Weakness Enumerations (CWEs) we have exposed in popular microprocessors. With our assistance, participants will get a real-world demonstration of the hardware fuzzing techniques used to detect these vulnerabilities and pinpoint their location in hardware design. We will explore the recent advances in hardware fuzzing using AI techniques and formal verification. We will use concrete, hands-on examples to quantitatively analyze the potential of these techniques for hardware security and the open challenges in the domain.

Technical Program – Oct. 27

10:00 - 11:30

Hardware Assurance in the Adversarial Era: Camouflage, Detection, and Lightweight Defense

Room: Ballroom A

Session Chair(s): Debdeep Mukhopadhyay, Indian Institute of Technology Kharagpur
Elif Bilge Kavun, Technische Universität Dresden

This session explores emerging techniques for securing hardware against adversarial threats. Topics include ML-driven camouflaging, RL-based Trojan detection, X-ray fingerprinting for trust verification, and lightweight defenses against bit-flip and inference attacks. Also featured are efficient hardware accelerators for post-quantum cryptography. Together, these works advance the state of resilient and trustworthy hardware design.

10:00

395: Designing with Deception: ML- and Covert Gate-Enhanced Camouflaging to Thwart IC Reverse Engineering

Junling Fan, David Koblah, Domenic Forte
University of Florida, United States

10:15

771: Defense in the Reverse Fragment: RL-Based Partial Netlist Hardware Trojan Detection

Bolun Li, Chen Dong, Decheng Qiu, Mingzhi Chen, Yang Yang
Fuzhou University, China

10:30

1043: TREX-F: TRustability of Electronics using X-ray based Fingerprinting

Tishya Sarma Sarkar {1}, Shuvodip Maitra {2}, Abhishek Chakraborty {2}, Sarani Bhattacharya {2}, Debdeep Mukhopadhyay {3}
{1} IIT, Kharagpur, India; {2} Indian Institute of Technology Kharagpur, India; {3} Department of Computer Science and Engineering, Indian Institute of Technology Kharagpur, India

10:45

1148: MIRAGE: Microarchitectural Footprints for Detecting Adversarial Attacks in One-Shot Inference

Soumi Chatterjee {1}, Debadrita Talapatra {1}, Nimish Mishra {1}, Aritra Hazra {2}, Debdeep Mukhopadhyay {3}
{1} Indian Institute of Technology Kharagpur, India; {2} Dept of CSE, IIT Kharagpur, India; {3} Department of Computer Science and Engineering, Indian Institute of Technology Kharagpur, India

Technical Program – Oct. 27

11:00

772: LEAF: Lightweight and Efficient Hardware Accelerator for Signature Verification of FALCON

Samuel Coulon {1}, Jinjun Xiong {2}, Jiafeng Xie {1}

{1} Villanova University, United States; {2} University at Buffalo, United States

11:15

1245: Non-Negative AdderNet: Algorithm-Hardware Co-design for Lightweight Defense of Bit-Flip Attacks

Yunxiang Zhang {1}, Sabbir Ahmed {2}, Abeer Almalky {2}, Adnan Siraj Rakin {3}, Wenfeng Zhao {3}

{1} Binghamton university, United States; {2} Binghamton University (SUNY), United States; {3} Binghamton University, United States

10:00 - 11:30

Journey from Efficient Layout to Successful Wafer

Room: Calgary

Session Chair(s): Ing-Chao Lin, National Yang Ming Chiao Tung University

Yuzhe Ma, The Hong Kong University of Science and Technology

(Guangzhou)

As the complexity of VLSI design has soared, it has become very challenging to ensure that the designs can actually be manufactured with high yield and reliability. This session will provide valuable insights into the intricate steps that contribute to the successful transition from VLSI layout to a wafer. We will delve into the critical aspects of layout efficiency, exploring various techniques and methodologies that ensure optimal performance and reliability. Additionally, we will discuss lithography simulation framework and advanced techniques for yield analysis and optimization.

10:00

832: GeoFA: A Geometric Finite Automaton Engine for Efficient Layout Pattern Matching

qingsheng qiu {1}, ziwen zheng {2}, boyu shi {1}, chao wang {1}

{1} southeast university, China; {2} southeast university

10:15

1179: G-Contour: GPU Accelerated Contour Tracing For Large-Scale Layouts

Shuo Yin {1}, Jiahao Xu {1}, Jiayi Jiang {1}, Mingjun Li {1}, Yuzhe Ma {2}, Tsung-Yi Ho {1}, Bei Yu {1}

{1} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} The Hong Kong University of Science and Technology (Guangzhou), China

10:30

432: 3D DRC: Design Rule Checking for 3D IC with U-Net-based Non-Manhattan Optimization

Shunjie Chang, Youran Wu, Jianli Chen, Jun Yu, Kun Wang

Fudan University, China

Technical Program – Oct. 27

10:45

554: LMLitho: A Large Vision Model-Driven Lithography Simulation Framework

Zhen Wang {1}, Hongquan He {1}, Tao Wu {1}, Xuming He {2}, Qi Sun {3}, Cheng Zhuo {3}, Bei Yu {4}, Jingyi Yu {1}, Hao Geng {1}
{1} ShanghaiTech University, China; {2} ShanghaiTech Univeristy, China; {3} Zhejiang University, China; {4} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China

11:00

589: CorDBA: Corners Decoupled Bayesian Approach for yield optimization

Yue Zhang {1}, Yunqi Li {2}, Shichang Ye {2}, Bojun Zhang {2}, Jinkai Wang {2}, Zhizhong Zhang {2}, Peng Wang {2}
{1} Beihang university, China; {2} Beihang University, China

11:15

354: BAGNet: A Boundary-Aware Graph Neural Network for SRAM Yield Analysis in Post-Layout Simulation

Haoyang Sang {1}, Changhao Yan {2}, Zhaori Bi {1}, Keren Zhu {1}, Xuan Zeng {1}
{1} Fudan University, China; {2} Associate Prof. Fudan University, China

10:00 - 11:30

Memory and Storage Design Optimization Techniques

Room: Ballroom B

Session Chair(s): Christian Pilato, Politecnico di Milano
Christian Haubelt, University of Rostock, Germany

This session showcases some cutting-edge system-level optimization techniques on memory and storage systems, including cache, CXL Memory, HBM, SSD, as well as near- and in-memory processing.

10:00

152: Rhea: a Framework for Fast Design and Validation of RTL Cache-Coherent Memory Subsystems

Davide Zoni, Andrea Galimberti, Adriano Guarisco
Politecnico di Milano, Italy

10:15

196: LsCMM-H: A TCO-Optimized Hybrid CXL Memory Expansion Architecture with Log Structure

Xingyu Chen {1}, Xiangrui Zhang {1}, Sirui Peng {1}, Zhiwang Guo {2}, Haidong Tian {3}, Xiankui Xiong {3}, Xiaoyong Xue {1}, Xiaoyang Zeng {1}
{1} State Key Laboratory of Integrated Chips and Systems, School of Microelectronics, Fudan University, China; {2} the State Key Laboratory of Integrated Chips and Systems, School of Microelectronics, Fudan University, China; {3} State Key Laboratory of Mobile Network and Mobile Multimedia Technology, ZTE Corporation, China

Technical Program – Oct. 27

10:30

229: STAR: Improving Lifetime and Performance of High-Capacity Modern SSDs Using State-Aware Randomizer

Omin Kwon {1}, Kyungjun Oh {1}, Jaeyong Lee {1}, Myungsuk Kim {2}, Jihong Kim {1}
{1} Seoul National University, Republic of Korea; {2} Kyungpook National University, Republic of Korea

10:45

705: Addressing Thermal Throttling in HBM

Gaurav Kothari, Kanad Ghose
Dept. of Computer Science, SUNY-Binghamton, United States

11:00

937: HD-MoE: Hybrid and Dynamic Parallelism for Mixture-of-Expert LLMs with 3D Near-Memory Processing

Haochen Huang {1}, Shuzhang Zhong {1}, Zhe Zhang {2}, Shuangchen Li {3}, Dimin Niu {4}, Hongzhong Zheng {5}, Runsheng Wang {1}, Meng Li {6}
{1} Peking University, China; {2} Alibaba, China; {3} Tsinghua University, China; {4} Alibaba Group, United States; {5} Damo Academy & Hupan Lab, China; {6} Institute for Artificial Intelligence and School of Integrated Circuits, Peking University, China

11:15

1339: CIMWise: An IREE-based End-to-End AI Compiler with Auto-Tuning for CIM Processors

Bo Mai {1}, Jin Wang {1}, Zhen Zhai {1}, Liang Zhang {1}, Yufu Zhang {2}, Longyang Lin {1}
{1} Southern University of Science and Technology, China; {2} Southern University of Science and Technology

Technical Program – Oct. 27

10:00 - 11:30

Speeding and Learning for Achieving 2D and 3D Design Closure

Room: Barcelona

Session Chair(s): Hung-Ming Chen, National Yang Ming Chiao Tung University
Jinjun Xiong, University of Buffalo

How can we speed up the design flow in modern 2D and 3D designs? In this session, we present how we use GPU, LLM, and unified flow tweaks towards better methodologies for design closure.

10:00

422: GPU Acceleration for Versatile Buffer Insertion

Yuan Pu {1}, Yuhao Ji {2}, Siying Yu {3}, Zuodong Zhang {4}, Zizheng Guo {5}, Zhuolun He {1}, Yibo Lin {5}, David Z. Pan {6}, Bei Yu {1}

{1} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {3} University of Illinois at Urbana-Champaign, United States; {4} School of Integrated Circuits, Peking University, China; {5} Peking University, China; {6} University of Texas at Austin, United States

10:15

1280: BUFFALO: PPA-Configurable, LLM-based Buffer Tree Generation via Group Relative Policy Optimization

Hao-Hsiang Hsiao {1}, Yi-Chen Lu {2}, Sung Kyu Lim {3}, Haoxing Ren {4}

{1} Georgia Institute of Technology, United States; {2} Nvidia, United States; {3} Georgia Tech, United States; {4} NVIDIA Corporation, United States

10:30

962: A Parallel Analytical Legalization Algorithm via Alternating Direction Method of Multipliers

Jaekyung Im, Seokhyeong Kang

Pohang University of Science and Technology, Republic of Korea

10:45

1091: A Unified Design Flow for Homogeneous and Heterogeneous 3D Integration with Fine-Pitch Hybrid Bonding

Gyumin Kim {1}, Heechun Park {2}

{1} Ulsan National Institute of Science & Technology, Republic of Korea; {2} Ulsan National Institute of Science and Technology (UNIST), Republic of Korea

11:00

1229: Snake-3D: Differentiable Learning for Cross-Tier Logic Path Snaking Optimization in 3D ICs

Yen-Hsiang Huang {1}, Sung Kyu Lim {2}

{1} Georgia Institute of Technology, United States; {2} Georgia Tech, United States

Technical Program – Oct. 27

11:15

479: SubtreeLU: High-Performance Parallel Sparse LU Factorization for Circuit Simulation

Jiawen Cheng, Yibin Zhang, Wenjian Yu

Tsinghua University, China

10:00 - 11:30

Special Session: Agile and Open Hardware Design and Verification

Room: Atlanta

Session Chair(s): Yun (Eric) Liang, Peking University

Compared to software design, hardware design is more expensive and time-consuming. The software development flow is agile partly because software engineers can leverage a rich set of open source tools such as runtime, middleware, etc., to get projects started and iterated easily and quickly. On the hardware side, the agile and open hardware ecosystem is rising thanks to the open ISA RISC-V and other emerging open EDA and IP tools. This not only saves costs, widens the design scope of hardware, but also provides a great opportunity for hardware innovation and creativity by designing and verifying custom hardware for different needs. This special session is formulated to be as broadly interesting and useful as possible to students, researchers and faculty, and to practice engineering in both EDA and hardware design area. This session is 2 hours, consisting of five talks. The first talk presents a machine learning accelerator design as RISC-V extension using an open source flow. The second talk develops an open CGRA framework. The third talk presents an FPGA-based open source RISC-V emulation platform. The fourth talk presents an agile hardware testing framework based on reinforcement learning. The last talk introduces an open-Source ASIP design framework.

10:00

20034: Invited Paper: Design of Machine Learning Accelerators as RISC-V Extensions using an Open Source Tool Flow

Batuhan Sesli, Muhammad Sabih, Frank Hannig, Jürgen Teich

Friedrich-Alexander-Universität Erlangen-Nürnberg (FAU), Germany

10:15

20035: Invited Paper: Building an Open CGRA Ecosystem for Agile Innovation

Rohan Juneja {1}, Pranav Dangi {1}, Thilini Kaushalya Bandara {2}, Zhaoying Li {1}, Dhananjaya Wijerathne {3}, Li-Shiuan Peh {1}, Tulika Mitra {1}

{1} National University of Singapore, Singapore; {2} Renesas Electronics; {3} Advanced Micro Devices

10:30

20036: Invited Paper: FEMU: An FPGA-Based Open-Source RISC-V Emulation Platform for Edge AI Systems Prototyping

Simone Machetti {1}, Deniz Kasap {1}, Juan Sapriza {1}, Rubén Rodríguez Álvarez {1}, Hossein Taji {1}, José Miranda {2}, Miguel Peón-Quirós {3}, David Atienza {1}

{1} École Polytechnique Fédérale de Lausanne (EPFL), Switzerland; {2} Center of Industrial Electronics (CEI), UPM, Madrid, Spain; {3} EcoCloud, EPFL, Lausanne, Switzerland

Technical Program – Oct. 27

10:45

20037: Invited Paper: CURE-Fuzz: Curiosity-Driven Reinforcement Learning for Enhanced Hardware Fuzzing in Agile Testing

Hanwei FAN {1}{2}, Ya Wang {1}, Xiaofeng Zhou {1}, Sicheng Li {3}, Binguang Zhao {3}, Yangdi Lyu {4}, Jiang Xu {4}, Wei Zhang {1}{2}

{1} Hong Kong University of Science and Technology, China; {2} Guangzhou HKUST Fok Ying Tung Research Institute; {3} Damo Academy & Hupan Lab, Alibaba, {4} The Hong Kong University of Science and Technology (Guangzhou)

11:00

20038: Invited Paper: APS: Open-Source Hardware-Software Co-Design Framework for Agile Processor Specialization

Youwei Xiao, Yuyang Zou, Yansong Xu, Yuhao Luo, Yitian Sun, Chenyun Yin, Ruifan Xu, Renze Chen, Yun Liang
Peking University, China

11:30 – 13:00

CEDA Luncheon Panel

Room: Ballroom A

Session Chair(s): Robert Wille, TUM, MQSC, SCCH

L. Miguel Silveira, IEEE CEDA President

13:00 – 14:30

Tutorial: Quantum Machine Learning Foundations for Quantum Architecture Design and Future Quantum EDA

Room: Sydney

Technical Program – Oct. 27

13:00 - 14:30

Hardware Reliability and Testing

Room: Ballroom B

Session Chair(s): Linda Milor, Georgia Institute of Technology
Zhiang Wang, Fudan University

This session introduces advanced solutions for robust, scalable, and efficient testing and reliability analysis of modern hardware systems, addressing challenges across compute-in-memory, RTL fault simulation, analog testing, and hardware fuzzing. It opens with CIMTester, which introduces a flexible BIST compiler tailored to compute-in-memory architectures. This is followed by RIROS, which proposes a two-dimensional parallelism framework that accelerates RTL fault simulation through unified scheduling and dynamic load balancing. The session continues with Tenpura, which achieves fast fault evaluation by narrowing the transient fault scope through scan chain-based fault analysis and leveraging FPGA-based emulation. Shifting focus to arithmetic resilience, SERA-Float presents a soft error–resilient floating-point format that avoids exceptions and bit-flip failures during deep learning inference with minimal impact on accuracy. This is followed by a novel adaptive testing method for SAR ADCs, which uses uncertainty-guided dynamic selection of measurement points. The session concludes with PROFUZZ, a directed graybox fuzzing framework that leverages ATPG’s structural analysis capabilities for effective verification of large hardware designs.

13:00

608: CIMTester: An Agile Golden-Result-Free BIST Compiler for Robust Compute-in-Memory

Wenjie Ren, Meng Wu, Mingxuan Li, Peiyu Chen, Tianyu Jia, Le Ye
Peking University, China

13:15

114: RIROS: A Parallel RTL Fault Simulation FFramework with Two-Dimensional Parallelism and Unified Schedule

Jiaping Tang {1}, Jianan Mu {2}, Zizhen Liu {3}, Ge Yu {4}, Tenghui Hua {4}, Bin Sun {4}, Silin Liu {4}, Jing Ye {1}, Huawei Li {1}

{1} State Key Lab of Processors, Institute of Computing Technology, Chinese Academy of Sciences/ University of Chinese Academy of Sciences/ CASTEST, China, China; {2} ICT, CAS, China; {3} State Key Lab of Processors, Institute of Computing Technology, Chinese Academy of Sciences, China; {4} State Key Lab of Processors, Institute of Computing Technology, Chinese Academy of Sciences/ University of Chinese Academy of Sciences, China

Technical Program – Oct. 27

13:30

530: Tenpura: A General Transient Fault Evaluation and Scope Narrowing Platform for Ultra-fast Reliability Analysis

Quan Cheng {1}, Huizi Zhang {2}, Chien-Hsing Liang {3}, Mingtao Zhang {1}, Jing-Jia Liou {3}, Jinjun Xiong {4}, Longyang Lin {2}, Masanori Hashimoto {1}
{1} Kyoto University, Japan; {2} Southern University of Science and Technology, China; {3} National Tsing Hua University, Taiwan; {4} University at Buffalo, United States

13:45

1109: SERA-Float: A Soft Error Resilient Approximate Floating-Point Computing Format

Vishesh Mishra {1}, Marcello Traiola {2}, Angeliki Kritikakou {3}, Olivier Sentieys {4}, Urbi Chatterjee {1}
{1} Indian Institute of Technology Kanpur, India; {2} Inria / IRISA, France; {3} Univ Rennes, Inria, CNRS, IRISA, France; {4} INRIA, France

14:00

1117: Uncertainty-Guided Live Measurement Sequencing for Fast SAR ADC Linearity Testing

Thorben Schey {1}, Khaled Karoonlatifi {2}, Michael Weyrich {1}, Andrey Morozov {1}
{1} University of Stuttgart, Germany; {2} Advantest Europe GmbH, Germany

14:15

973: PROFUZZ: Directed Graybox Fuzzing via Module Selection and ATPG-Guided Seed Generation

Raghul Saravanan {1}, Sudipta Paria {2}, Aritra Dasgupta {2}, SWARUP BHUNIA {2}, Sai Manoj Pudukotai Dinakarrao {1}
{1} George Mason University, United States; {2} University of Florida, United States

13:00 - 14:30

Logic Synthesis Reimagined: Mapping, Rewriting, and Structural Optimization

Room: Ballroom A

Session Chair(s): Siva Satyendra Sahoo, imec

Yun Liang, Peking University

This session showcases foundational advances in logic synthesis, covering structural rewrites, delay and area optimization, and technology mapping. Techniques include gradient descent, SAT solving, and e-graph-based transformations — all pushing the boundaries of what logic optimization tools can deliver.

13:00

79: ExactMap: Enhancing Delay Optimization in Parallel ASIC Technology Mapping

Zhenxuan Xie {1}, Lixin Liu {2}, Tianji Liu {1}, Evangeline Young {1}
{1} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} The Chinese University of Hong Kong, China

Technical Program – Oct. 27

13:15

303: GradMap: A Gradient-Descent Approach to Simultaneous Technology Mapping, Buffer Insertion, and Gate Sizing

HsinYing Tsai, Chung-Kai Wu, Chih-Cheng Hsu, Jie-Hong Roland Jiang
National Taiwan University, Taiwan

13:30

301: Versatile Rewiring and Concurrent Resynthesis for High-Quality Customized Optimization

JiunHao Chen {1}, Jie-Hong Roland Jiang {1}, Alan Mishchenko {2}
{1} National Taiwan University, Taiwan; {2} UC Berkeley, United States

13:45

1293: Revisit Choice Network for Synthesis and Technology Mapping

Chen Chen, Jiaqi Yin, Cunxi Yu
University of Maryland, College Park, United States

14:00

341: MuSTNet: SAT-based Exact Multi-Stage Transistor Network Synthesis with Placement Awareness

Jiun-Cheng Tsai {1}, Wei-Min Hsu {1}, Kuei-Lin Wu {1}, Hsuan-Ming Huang {1}, Jen-Hang Yang {1}, Heng-Liang Huang {1}, Yen-Ju Su {2}, Hung-Pin (Charles) Wen {2}
{1} Mediatek, Taiwan; {2} National Yang Ming Chiao Tung University, Taiwan

14:15

400: EqMap: FPGA LUT Remapping using E-Graphs

Matthew Hofmann, Berk Gokmen, Zhiru Zhang
Cornell University, United States

Technical Program – Oct. 27

13:00 - 14:30

Processing-in-Memory (PIM/CIM) Powered AI Acceleration

Room: Barcelona

Session Chair(s): Hiromitsu Awano, Kyoto University

Yufei Ma, Peking University

This session highlights Processing-in-Memory (PIM) as a key technology to overcome the memory wall in AI. Papers showcase the evolution of digital CIM with architectures like Adder-DCIM (eliminating multipliers) and the reconfigurable PAR-CIM. The focus then shifts to applying PIM to LLMs, where LLM-on-the-Palm uses PIM-enhanced Flash, LEAP combines PIM with a NoC for balanced dataflow, and LP-Spec optimizes for speculative inference. Addressing practical hardware challenges, OptiRange presents a method to optimize costly ADCs in ReRAM-based accelerators.

13:00

1: Adder-DCIM: A Parallel Bit-Flexible Digital CIM Joint Model Compression Framework for AdderNet Inference

Haikang Diao {1}, Chuyue Tang {1}, Bocheng Xu {2}, Haoyang Luo {1}, Meng Li {3}, Yuan Wang {4}, Xiyuan Tang {1}

{1} Peking University, China; {2} Peking university, China; {3} Institute for Artificial Intelligence and School of Integrated Circuits, Peking University, China; {4} Peking University, China

13:15

8: PAR-CIM: A Precise/Approximate Reconfigurable Digital CIM Macro with 0.35-4b Fractional Mixed-Bitwidth Quantization

Han Zhang {1}, Zhenyu Xue {1}, Wente Yi {1}, Tianshuo Bai {1}, Lehao Tan {1}, Jingcheng Gu {1}, Weijie Ding {1}, Wang Kang {1}, Biao Pan {2}

{1} Beihang University, China; {2} Beihang University, China

13:30

247: LLM-on-the-Palm: Mobile LLM Inference with PIM-Enhanced NAND Flash Memory

Hyunjin Kim, Sanghyeok Han, Jae-Joon Kim

Seoul National University, Republic of Korea

13:45

323: LEAP: LLM Inference on Scalable PIM-NoC Architecture with Balanced Dataflow and Fine-Grained Parallelism

Yimin Wang, Yue Jiet Chong, Xuanyao Fong

National University of Singapore, Singapore

14:00

394: LP-Spec: Leveraging LPDDR PIM for Efficient LLM Mobile Speculative Inference with Architecture-Dataflow Co-Optimization

Siyuan He {1}, Zhantong Zhu {2}, Yandong He {2}, Tianyu Jia {2}

{1} Beijing, China; {2} Peking University, China

Technical Program – Oct. 27

14:15

948: OptiRange: An Efficient ReRAM-Based PIM Accelerator with ADC Resolution Optimization

Sangkyu Jeon {1}, Gisan Ji {2}, Yeonggeon Kim {3}, Youngjun Park {2}, Sangyeon Kim {2}, Sungju Ryu {2}

{1} Sogang university, Republic of Korea; {2} Sogang University, Republic of Korea; {3} sogang university, Republic of Korea

13:00 - 14:30

Silicon Whispers: Side Channels, Rowhammer, and Microarchitectural Leaks

Room: Calgary

Session Chair(s): Jiafeng Xie, Villanova University

Akash Kumar, Ruhr University Bochum

This session explores subtle "whispers" of the silicon, from transient execution and cache-timing attacks to physical side channels like power and jitter analysis. The session covers a wide spectrum, demonstrating vulnerabilities and defenses in systems ranging from post-quantum cryptography to AI accelerators.

13:00

989: SpectrePrefetch: Undermining Cache-Centric Secure Speculation with Modern Hardware Prefetchers

Fang Jiang {1}, Fei Tong {1}, Xiaoyu Cheng {1}, Zhe Zhou {1}, Hongyu Wang {2}, Yuxing Mao {3}

{1} Southeast University, China; {2} Nanjing Unipower Information Technology Co., Ltd, China; {3} Chongqing University, China

13:15

1436: CacheGuardian: A Timing Side-Channel Resilient LLC Design

Ziang Zhou {1}, Qi Zhu {1}, Hao Lan {1}, Huifeng Zhu {2}, Wei Yan {3}, Chenglu Jin {4}, xuejun an {5}, Xiaochun Ye {5}

{1} Institute of Computing Technology, Chinese Academy of Science; University of Chinese Academy of Sciences, Beijing, China, China; {2} Washington University in St.Louis, United States; {3} Institute of ComputingTechnology, Chinese Academy of Sciences, China; {4} CWI Amsterdam, Netherlands; {5} Institue of computing technology, CAS, China

13:30

1059: Everything Depends on Your Hammer: A Systematic Rowhammer Attack Exploration on SPHINCS+

S. G. Shoaib Ahamed, Mrityunjay Shukla, Khushang Singla, Sayandeep Saha
Indian Institute of Technology Bombay, India

Technical Program – Oct. 27

13:45

1008: The Fellowship of the Leak: Power Analysis of a Masked FrodoKEM Hardware Accelerator

Martin Schmid {1}, Giuseppe Manzoni {1}, Aydin Aysu {2}, Elif Bilge Kavun {3}
{1} University of Passau, Germany; {2} North Carolina State University, United States; {3} Barkhausen Institut & TU Dresden, Germany

14:00

1095: Leaks beyond Bits: Deep Learning-Assisted Side-Channel Attacks on Hyperdimensional Computing Accelerators

Brojogopal Sapui {1}, Mehdi Tahoori {2}
{1} Karlsruhe Institute of Technology, Germany, Germany; {2} Karlsruhe Institute of Technology, Germany

14:15

97: JitFilt: Mitigate Jitter-based Side Channel Analysis Attacks

Darshana Jayasinghe {1}, Yuanhua Zhong {2}, Sri Parameswaran {2}
{1} University of Sydney, Australia; {2} The University of Sydney, Australia

13:00 - 14:30

Special Session: Hardware-Software Co-Design for highly Optimized, Customized, and Reliable AI Systems

Room: Atlanta

Session Chair(s): Joerg Henkel, Karlsruhe Institute of Technology
Mehdi Tahoori, Karlsruhe Institute of Technology

Conventional HW-SW co-design approaches using C/C++, SystemC, and standard co-simulation tools are inadequate for modern AI systems, where hardware and software are deeply interdependent. A unified co-design loop is needed to jointly optimize performance, energy, and reliability, especially for embedded AI under tight resource constraints. Key challenges include fragmented toolchains, discrete hardware behavior within continuous AI optimization, and limited scalability of co-simulation. Reliability is critical, with silent data corruption (SDC) from both hardware faults and software errors posing significant risks. This session explores co-design strategies for building robust, energy-efficient, and reconfigurable AI architectures. Emphasis will be placed on practical methods to close the HW-SW gap and support scalable, reliable AI deployment.

13:00

Driving the Next Wave of Efficiency: Hardware-Software Co-Design for AI Systems

Joerg Henkel, Mehdi Tahoori
Karlsruhe Institute of Technology, Germany

Technical Program – Oct. 27

13:15

A3C3 – AI Algorithm & Accelerator Co-design, Co-search, and Co-generation

Selin Yildirim

University of Illinois Urbana-Champaign, United States

13:30

SDC Mitigation in AI Workloads running on Scale-Out Systems with Large Mission Time

Nirmal R. Saxena

NVIDIA, United States

13:45

Modalix - A HW/SW co-designed multi-modal MLSoC for GenAI and Smart Vision

Srivi Dhruvanarayan

SIMA, United States

14:30 – 15:00

Coffee Break

Room: Foyer Ballroom

15:00 - 16:30

Advanced Modeling and Optimization for Power and Thermal Integrity

Room: Ballroom B

Session Chair(s): Masanori Hashimoto, Kyoto University

Hai Wang, University of Electronic Science and Technology of China

The power integrity and thermal concern are critical issues in nowadays IC design. This session assembles the advanced modeling, analysis and optimization techniques for power and thermal integrity. The first two papers address the optimization problems in decoupling capacitor placement and the placement of voltage regulators in 2.5D ICs, respectively. The third and fourth papers present the approaches on dynamic power estimation and the PDN circuit reduction, both leveraging cutting-edge deep learning technologies. These techniques are useful for PDN design and power integrity analysis. The last two papers present an electro-thermal co-analysis method based on PINN and a pre-layout parasitic estimation technique based on graph contrastive learning, respectively.

15:00

582: Semidefinite Programming-Based Decoupling Capacitor Placement for Power Distribution Network Optimization

Zong-Ying Cai {1}, Wei-Han Mao {1}, Yao-Wen Chang {1}, Yang Lu {2}, Jerry Bai {2}, Bin-Chyi Tseng {2}

{1} National Taiwan University, Taiwan; {2} ASUSTeK Computer Inc., Taiwan

Technical Program – Oct. 27

15:15

494: Optimal Selection and Placement of Voltage Regulators in 2.5D Heterogeneously Integrated Systems

Hangyu Zhang {1}, Divya Yogi {2}, Sachin S. Sapatnekar {2}, Ramesh Harjani {2}

{1} University of Minnesota Twin Cities, United States; {2} University of Minnesota, United States

15:30

431: VIRTUAL: Vector-based Dynamic Power Estimation via Decoupled Multi-Modality Learning

Yuntao Lu {1}, Mingjun Wang {2}, Yihan Wen {3}, Boyu Han {4}, Jianan Mu {5}, Huawei Li {6}, Bei Yu {7}

{1} The Chinese University of Hong Kong, China; {2} State Key Lab of Processors, Institute of Computing Technology, Chinese Academy of Sciences; University of Chinese Academy of Sciences; CASTEST, Beijing, China; {3} Beijing University of Technology, China; {4} Stanford, United States; {5} ICT, CAS, China; {6} Institute of Computing Technology, Chinese Academy of Sciences, China; {7} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China

15:45

894: MF-MOR: Multi-Fidelity Model Order Reduction for Many-Port Linear Systems in Chip Power Modeling

Zhenjie Lu {1}, Hang Zhou {2}, Quan Chen {2}

{1} South University of Science and Technology, China; {2} School of Microelectronics, Southern University of Science and Technology, China

16:00

941: Accelerating Electro-Thermal Co?Analysis via Coarse?to?Fine Physics?Informed Neural Networks

Xiao Dong, Songyu Sun, Xunzhao Yin, Zhou Jin, Zhiguo Shi, Cheng Zhuo
Zhejiang University, China

16:15

641: Transferable Parasitic Estimation via Graph Contrastive Learning and Label Rebalancing in AMS Circuits

Shan Shen {1}, Shenglu Hua {2}, Jiajun Zou {1}, Jiawei Liu {2}, Jianwang Zhai {2}, Chuan Shi {3}, Wenjian Yu {4}

{1} Nanjing University of Science and Technology, China; {2} Beijing University of Posts and Telecommunications, China; {3} School of Computer Science, Beijing University of Posts and Telecommunications, China; {4} Tsinghua University, China

Technical Program – Oct. 27

15:00 - 16:30

Edge Intelligence: Co-Design, Resilience, and Efficiency

Room: Barcelona

Session Chair(s): Qi Zhu, Northwestern University
Cheng Zhuo, Zhejiang University

This session showcases cutting-edge innovations in edge computing that span algorithm-architecture co-design, resource-aware scheduling, and resilient systems for intermittently powered or resource-constrained environments. From ultra-efficient graph analytics to robust checkpointing and fine-grained PCIe tracing, the selected works push the boundaries of edge intelligence design. The session also highlights system-level optimizations for mitigating resource contention and coordinating compound AI workloads across heterogeneous platforms.

15:00

1019: Scalable and Asynchronous Differential Checkpointing for Intermittently Powered Devices

Youngbin Kim, Yoojin Lim
ETRI, Republic of Korea

15:15

1503: Ultra Energy-Efficient Butterfly Counting in Bipartite Networks via Algorithm-Architecture Co-Optimization

Haiyang Liu {1}, Xueyan Wang {2}, Jianlei Yang {2}, Xiaotao Jia {3}, Gang Qu {4}, Weisheng Zhao {2}

{1} BEIHANG University, China; {2} Beihang University, China; {3} School of Integrated Circuit Science and Engineering, Beihang University, China; {4} Univ. of Maryland, College Park, United States

15:30

54: DevTrace: Efficient and Fine-grained PCIe Transaction Tracing for Edge Intelligence Workloads

Zhibai Huang {1}, Kailiang Xu {2}, Zhixiang Wei {2}, Yinghao Deng {2}, Chen Chen {2}, James Yen {3}, Yun Wang {2}, Fangxin Liu {3}, MingYuan Xia {4}, Zhengwei Qi {2}

{1} Shanghai Jiao tong university, China; {2} Shanghai Jiao Tong University, China; {3} Shanghai Jiaotong University, China; {4} UltraRISC, China

15:45

98: Mitigating Resource Contention for Responsive On-device Machine Learning Inferences

Minsung Kim {1}, Jihoon Lee {1}, Seongjin Chou {1}, Whisoo Chung {2}, Inwoo Kim {1}, Woosung Kang {3}, Hyosu Kim {4}, Sangeun Oh {5}, Hoon Sung Chwa {3}, Kilho Lee {1}

{1} Soongsil University, Republic of Korea; {2} Seoul National University, Republic of Korea; {3} DGIST, Republic of Korea; {4} Chung-Ang University, Republic of Korea; {5} Korea university, Republic of Korea

Technical Program – Oct. 27

16:00

1217: Twill: Scheduling Compound AI Systems on Heterogeneous Mobile Edge Platforms

Zain Taufique {1}, Aman Vyas {1}, Antonio Miele {2}, Pasi Liljeberg {1}, Anil Kanduri {1}
{1} University of Turku, Finland; {2} Politecnico di Milano, Italy

15:00 - 16:30

In memory computing and computing in memory

Room: Calgary

Session Chair(s): Martin Margala, University of Louisiana at Lafayette
Sarani Bhattacharya, IIT Kharagpur

Von Neumann computing models are increasingly facing the challenge of memory bandwidth due to growing data-intensive workloads. In memory computing, in its various forms, provides a promising direction to alleviate the bandwidth challenge. This session will dive into the latest and greatest in computing in memory and in memory computing using both classical and emerging memory technologies.

15:00

235: HyDra: SOT-CAM Based Vector Symbolic Macro for Hyperdimensional Computing

Md Mizanur Rahaman Nayan, Che-Kai Liu, Zishen Wan, Arijit Raychowdhury, Azad Naeemi
Georgia Institute of Technology, United States

15:15

33: Row-Column Hybrid Grouping for Fault-Resilient Multi-Bit Weight Representation on IMC Arrays

Kang Eun Jeon {1}, Sang Heum Yeon {1}, Jinhee Kim {1}, Hyeonsu Bang {1}, Johnny Rhe {1}, Jong Hwan Ko {2}
{1} Sungkyunkwan University, Republic of Korea; {2} Sungkyunkwan University (SKKU), Republic of Korea

15:30

621: DANCE: Dual-Side Agile N:M Sparse Compressed Digital CiM Accelerator for Efficient Compound AI

Zhonghao Chen, Hongtao Zhong, Jianhe Deng, Mulin Shi, Yongpan Liu, Huazhong Yang, Xueqing Li
Tsinghua University, China

15:45

1240: Energy-efficient Multi-Operand XOR Logic-based CIM Accelerator using RRAM technology

Abhairaj Singh {1}, Konstantinos Stavrakakis {1}, Rajendra Bishnoi {2}, Rajiv Joshi {3}, Said Hamdioui {4}
{1} TU Delft, Netherlands; {2} Delft University of Technology, Netherlands; {3} IBM, United States; {4} Delft University of Technology, Netherlands

Technical Program – Oct. 27

16:00

1022: An Adaptive Sparse Matrix Compression CIM Accelerator based on 256Kb SOT-MRAM for Downlink Massive MIMO Communications

Liangchen Li {1}, Jianxin Wu {1}, Changyu Li {1}, Liang Zhang {1}, Anyang Yu {1}, Junda Zhao {1}, Zhaohao Wang {1}, Chengyuan Sun {2}, Kaihua Cao {1}, Hongxi LIU {2}, Wang Kang {1}, He Zhang {1}, Weisheng Zhao {1}

{1} Beihang University, China; {2} Truth Memory Corporation, China

16:15

1247: JSA-CIM: A Joint-Sparse AdderNet Compute-In-Memory Accelerator for Energy-Efficient Edge AI Applications

Omar Al Kailani {1}, Yunxiang Zhang {2}, Wenfeng Zhao {3}

{1} SUNY Binghamton, United States; {2} Binghamton university, United States; {3} Binghamton University, United States

15:00 - 16:30

Power and Performance Conscious Hardware Design

Room: Ballroom A

Session Chair(s): Georgios Zervakis, University of Patras, Greece
Tianyu Jia, Peking University

Power and energy efficiency have become a key consideration across all layers of hardware design. This session brings together a set of approaches that aim to improve power estimation, reduce energy consumption, and support efficient deployment of modern workloads. Topics include early power modeling directly in high-level synthesis, circuit-architecture co-design of voltage-stacked processors for power delivery network efficiency, and scheduling techniques for logic-in-memory systems that minimize data movement. The session also covers GNNs for rapid and transferable design space exploration in approximate computing, as well as lookup-based solutions for deploying learned activation functions on FPGAs. Together, these works explore methods that balance power, performance, and flexibility across various design stages and application domains.

15:00

821: TickTockStack: In-Datapath Current Imbalance Elimination Using Clocked Differential Logic in a Voltage Stacked Vector Processor

Michal Gorywoda, Wanyeong Jung

KAIST, Republic of Korea

15:15

1419: MASIM: An Energy-Efficient Multi-Array Scheduler for SIMD Logic-in-Memory Architectures

Xingyue Qian {1}, Chen Nie {1}, Zhezhi He {2}, Weikang Qian {1}

{1} Shanghai Jiao Tong University, China; {2} School of Electronic Information and Electrical Engineering, Shanghai Jiao Tong University, China

Technical Program – Oct. 27

15:30

586: HIPPO: A Hierarchy-Preserving and Noise-Tolerant Pre-HLS Power Modeling Framework for FPGA

Zefan Lin {1}, Zedong Peng {2}, Mingzhe Gao {2}, Jieru Zhao {2}, Zhe Lin {1}
{1} Sun Yat-sen University, China; {2} Shanghai Jiao Tong University, China

15:45

691: ApproxGNN: A Pretrained GNN for Parameter Prediction in Design Space Exploration for Approximate Computing

Ondrej Vlcek, Vojtech Mrazek
Brno University of Technology, Czech Republic

16:00

839: Optimizing Neural Networks with Learnable Non-Linear Activation Functions via Lookup-Based FPGA Acceleration

Mengyuan Yin {1}, Benjamin Choong {1}, Chuping Qu {1}, Rick Goh {1}, Weng-Fai Wong {2}, Tao Luo {1}
{1} Agency for Science, Technology and Research, Singapore; {2} National University of Singapore, Singapore

15:00 - 16:30

Special Session: EDA from GPU Acceleration to GPU Exploration

Room: Sydney

Session Chair(s): Wen-Hao Liu, NVIDIA

In recent years, many GPU-accelerated EDA algorithms have been introduced, demonstrating runtime speedups of hundreds to thousands of times for specific stages in the EDA flow. Beyond runtime improvements, this special session will explore various GPU adoption models in the EDA flow that also enhance design quality, particularly in terms of power and area reduction. The session will feature a total of five talks. Three of these are invited presentations from industry leaders—Cadence, Synopsys, and Siemens—who will showcase their production-ready, GPU-based EDA solutions. The remaining two talks are from academic research groups, offering insights into emerging GPU adoption models with the potential to further improve EDA design quality.

15:00

Unleashing Performance: GPU-Accelerated Multi-Physics Simulation for Next-Generation Engineering

Albert Zeng
Cadence, United States

15:15

Unlocking New Levels of Efficiency in Semiconductor Design and Verification with AI and Accelerated Computing

Ahmed Ramadan
Synopsys, United States

Technical Program – Oct. 27

15:30

Accelerating the Future of Lithography: GPU-Optimized ILT for Full-Chip Implementation

Germain Fenger

Siemens EDA, United States

15:45

20001: Invited Paper: Beyond Speed: Leveraging GPU Exploration for Enhanced Solution Quality in VLSI Physical Design

Zhiang Wang

University of California San Diego, USA

16:00

20002: Invited Paper: Generalized GPU-Accelerated Dynamic Programming with Application to Mixed-Cell-Height Detailed Placement

Da-Wei Huang{1}, Shao-Yun Fang{1}, and Wen-Hao Liu{2}

{1} National Taiwan University of Science and Technology, Taiwan; {2} NVIDIA, Taiwan

15:00 - 16:30

Special Session: Security Challenges and Opportunities in In-Sensor Computing Systems

Room: Atlanta

Session Chair(s): Qiaoyan Yu, University of New Hampshire

Traditionally, massive sensor data are continuously collected to support external processing like Artificial Intelligence (AI)-based analysis and prediction, which consumes significant power and introduces delays. To address this limitation, we need a sensing system that is power-efficient, compact, and capable of processing large amounts of data in real time without sacrificing accuracy. One promising direction is to enable sensing units to offer additional computing capability than a simple sensor. In-Sensor Computing (ISC) emerges as a new computation paradigm to address the increasing concern on latency and energy consumption in sensory data transmission, analog-to-digital conversion (ADC), and data pre-processing. Existing literature for ISC mainly investigates the materials and device structure to implement desired functions, pursue better performance and examine the feasible integration technologies. While offering significant benefits in performance improvement and power saving, in-sensor computing could also bring new security challenges. Unfortunately, limited work is available to study new threats and unique attack surfaces for ISC systems. In this special session, four speakers will present the state-of-the-art studies in this research area and show their visions.

15:00

20028: Invited Paper: Security Under the Lens: Vulnerabilities in In-Sensor Computing Systems

Mashrafi Kajol {1}, Wei Lu {2}, Qiaoyan Yu {1}

{1} University of New Hampshire, United States; {2} Keene State College, United States

Technical Program – Oct. 27

15:15

20029: Invited Paper: Toward Secure In-Sensor Intelligence: Threats and Defenses in SNNs

Archisman Ghosh, Swaroop Ghosh
Pennsylvania State University, United States

15:30

20030: Invited Paper: Rowhammer Mitigation by Approximate Computing: A Compressed Sensing Case Study

Yuhang Hao, Yun Wu, Minmin Jiang, Maire O'Neill, Chongyan Gu
Queen's University Belfast, United Kingdom

15:45

20031: Invited Paper: Side Channel Vulnerability Analysis of Printed Neuromorphic Circuits

Priyanjana Pal, Brojogopal Sapui, Mehdi Tahoori
Karlsruhe Institute of Technology, Germany

16:30 – 16:45

Transition

16:45 - 17:45

Advancing towards the Optimum Cell Generation

Room: Barcelona

Session Chair(s): Bei Yu, CHUK

Josef Mittermaier, Cadence

There has been a long while when we optimize and automate the generation of cell libraries, not to mention difficult advanced nodes. In this session, we have numerous generators to optimize in simultaneous objectives.

16:45

1181: CoP&R: Co-Optimizing Place-and-Route for Standard Cell Layout via MCTS and AIISAT

Yen-Ju Su {1}, Jiun-Cheng Tsai {2}, Hsuan-Ming Huang {2}, Aaron C.-W. Liang {1}, Han-Ya Tsai {1}, Wei-Min Hsu {2}, Jen-Hang Yang {2}, Heng-Liang Huang {2}, Hung-Pin (Charles) Wen {1}

{1} National Yang Ming Chiao Tung University, Taiwan; {2} Mediatek, Taiwan

17:00

368: Synthesis of Standard Cells of Minimum Delay

Sehyeon Chung, Hyunbae Seo, Taewhan Kim
Seoul National University, Republic of Korea

Technical Program – Oct. 27

17:15

414: SO3-Cell: Standard Cell Layout Synthesis Framework for Simultaneous Optimization of Topology, Placement, and Routing

Chung-Kuan Cheng {1}, Andrew Kahng {1}, Byeonggon Kang {2}, Seokhyeong Kang {3}, Jakang Lee {3}, Bill Lin {1}

{1} UCSD, United States; {2} University of California San Diego, United States; {3} Pohang University of Science and Technology, Republic of Korea

17:30

1098: Standard Cell Layout Generator Empowered by ILP-based Routing with Dynamic Grid-Shifting for Advanced Nodes

Zhengzhe Zheng {1}, Keyu Peng {1}, YINUO Wu {2}, Hao Gu {1}, chao wang {3}, Ziran Zhu {2} {1} Southeast University, China; {2} School of Integrated Circuits, Southeast University, China; {3} teacher, China

16:45 - 17:45

Co-Design and Compression for Efficient AI at the Edge

Room: Sydney

Session Chair(s): Daniel Müller-Gritschneider, TU Wien

Marcel Walter, Technical University of Munich

This session focuses on innovative algorithm-hardware co-design strategies and compression methods for pushing neural network performance to the edge. Papers explore precision-aware training, temporal logic for model compression, and optimization frameworks that cater to low-power and latency-sensitive environments. Learn how hybrid architectures and efficient learning techniques are shaping the future of on-device intelligence.

16:45

44: Perturbation-efficient Zeroth-order Optimization for Hardware-friendly On-device Training

Qitao Tan {1}, Sung-En Chang {2}, Rui Xia {3}, Huidong Ji {4}, Chence Yang {1}, Ci Zhang {1}, Jun Liu {2}, zheng zhan {2}, Zhenman Fang {5}, zhuo zou {4}, Yanzhi Wang {2}, Jin Lu {1}, Geng Yuan {1}

{1} University of Georgia, United States; {2} Northeastern University, United States; {3} University of Pennsylvania, United States; {4} Fudan University, China; {5} Simon Fraser University, Canada

17:00

83: MiCo: End-to-End Mixed Precision Neural Network Co-Exploration Framework for Edge AI

Zijun JIANG {1}, Yangdi Lyu {2}

{1} Hong Kong University of Science & Technology (Guangzhou), China; {2} Hong Kong University of Science and Technology (Guangzhou), China

Technical Program – Oct. 27

17:15

286: TOGGLE: Temporal Logic-Guided Large Language Model Compression for Edge

Khurram Khalil, Khaza Anuarul Hoque

University of Missouri, United States

17:30

374: BARQ: Boundary-Aware Regularized Training for Accurate Inference on Computing-in-Memory Accelerators with Low-Precision A/D Conversion

Tingrui Ren {1}, Bi Wang {1}, Liang Wang {2}, Yuanfu Zhao {2}

{1} The School of Integrated Circuit Science and Engineering, Beihang University, China;

{2} The Beijing Microelectronics Technology Institute, China

16:45 - 17:45

Emerging Architectures and Automation for Neuromorphic, In-Memory, and Microfluidic Computing

Room: Calgary

Session Chair(s): Cindy Yang Yi, Virginia Tech

Bing Li, University of Siegen, Germany

This session brings together recent innovations at the intersection of hardware design and intelligent computation, highlighting new architectures and methodologies for accelerating non-traditional computing paradigms. The first paper introduces PLAIN, a processing-in-memory (PIM) accelerator that leverages high internal bandwidth and mixed-precision quantization to enhance large language model (LLM) inference. The second paper presents LUT-HD, a lightweight yet efficient architecture for hyperdimensional computing (HDC) inference based on optimized table lookup mechanisms. The third paper, NeuroPDE, explores a neuromorphic approach to solving partial differential equations using spintronic and ferroelectric devices, pushing the frontier of unconventional physics-based computing. The fourth paper proposes an automatic design framework for modular microfluidic routing blocks, streamlining the creation of reconfigurable and scalable biochip systems. Together, these works span a diverse range of emerging computing substrates, showcasing how architectural and design innovations can enable new levels of performance, functionality, and adaptability across domains.

16:45

445: PLAIN: Leveraging High Internal Bandwidth in PIM for Accelerating Large Language Model Inference via Mixed-Precision Quantization

Yiwei Hu {1}, Fangxin Liu {2}, Zongwu Wang {2}, Yilong Zhao {3}, Tao Yang {4}, Haibing Guan {3}, Li Jiang {3}

{1} Shanghai Jiaotong university, China; {2} Shanghai Jiaotong University, China; {3}

Shanghai Jiao Tong University, China; {4} Huawei Technologies Co., Ltd, China

Technical Program – Oct. 27

17:00

975: LUT-HD: Accelerating Hyperdimensional Computing Inference via Efficient Table Lookup

Haodong Lu {1}, Da Tang {1}, Xiqiong Bai {2}, Zexu Zhang {1}, Tianyi Zhou {1}, Xinran Li {1}, Kun Wang {1}

{1} Fudan University, China; {2} Nanjing University Of Posts And Telecommunications, China

17:15

130: NeuroPDE: A Neuromorphic PDE Solver Based on Spintronic and Ferroelectric Devices

Siqing Fu {1}, Lizhou Wu {2}, Tiejun Li {2}, Chunyuan Zhang {2}, Sheng ma {3}, Jianmin Zhang {2}, Yuhang Tang {3}, Jixuan Tang {2}

{1} College of Computer Science and Technology, National University of defense technology, China; {2} College of Computer Science and Technology, National University of Defense Technology, China; {3} National University of Defense Technology, China

17:30

1158: Automatic Design for Modular Microfluidic Routing Blocks

Philipp Ebner {1}, Maria Emmerich {2}, Eric Safai {3}, Aniruddha Paul {3}, Mathieu Odiijk {3}, Joshua Loessberg-Zahl {3}, Robert Wille {2}

{1} Johannes Kepler University, Austria; {2} Technical University of Munich, Germany; {3} University of Twente, Netherlands

16:45 - 17:45

Learning-Based Hardware Optimization and Scheduling

Room: Atlanta

Session Chair(s): Holger Froening, Heidelberg University

Focusing on adaptive optimization techniques, this session presents novel learning-based frameworks for circuit scheduling, routing, power management, and device co-design. The featured papers introduce reinforcement learning, coreset-based transfer learning, and physics-informed modeling to tackle multi-objective constraints across scales. Collectively, these approaches highlight the value of intelligent learning strategies in navigating complex hardware optimization landscapes efficiently and effectively.

16:45

448: M3: Mamba-assisted Multi-Circuit Optimization via Model-based RL with Effective Scheduling

Youngmin Oh {1}, Jinje Park {1}, Taejin Paik {2}, Seunggeun Kim {3}, Suwan Kim {4}, Yoon Hyeok Lee {2}, David Z. Pan {3}

{1} Samsung Advanced Institute of Technology, Republic of Korea; {2} Samsung Electronics, AI Center, Republic of Korea; {3} University of Texas at Austin, United States; {4} Samsung Electronics, Republic of Korea

Technical Program – Oct. 27

17:00

795: Towards Multi-Objective Routing: A Novel Coreset-based Transfer Learning Framework

Xianglu Wang, Hu Ding

University of Science and Technology of China, China

17:15

463: EarDVFS: Environment-Adaptable RL-based DVFS for Mobile Devices

Jaeheon Kwak {1}, Sangeun Oh {2}, Jinkyu Lee {3}, Insik Shin {1}

{1} KAIST, Republic of Korea; {2} Korea university, Republic of Korea; {3} Sungkyunkwan University (SKKU), Republic of Korea

17:30

704: PiGen: Accelerating Re-RAM Co-Design via Generative Physics-Informed Modeling

Zihan Zhang, Marco Donato

Tufts University, United States

16:45 - 17:45

Optimization and Simulation From Circuits to Systems

Room: Ballroom A

Session Chair(s): Quan CHEN, SUSTech, China

Wei Xing, University of Sheffield, UK

Explore revolutionary optimization paradigms that push the boundaries of computational efficiency and design quality. This session features breakthrough numerical methods for nonlinear circuit analysis, information-theoretic approaches to parameter optimization, and GPU-accelerated algorithms for large-scale photonic integration. From exponential integrators that outperform traditional methods to dynamic variable expansion mimicking expert designers' workflows, these innovations demonstrate how sophisticated optimization strategies are enabling unprecedented scalability and performance in modern circuit and system design.

16:45

596: EI-TR: A Versatile Exponential Integrator Framework for Transient Analysis of Generic Nonlinear Circuits

Hang Zhou {1}, Zhenjie Lu {2}, Quan Chen {1}

{1} School of Microelectronics, Southern University of Science and Technology, China; {2} South University of Science and Technology, China

17:00

1260: DIVE: Dynamic Information-Guided Variable Expansion for Deeper Analog Circuit Optimization

Zhuohua Liu {1}, Weilun Xie {2}, Yuxuan Zhang {1}, Chen Wang {1}, Yuanqi Hu {1}, Wei Xing {3}

{1} Beihang University, China; {2} Shenzhen University, China; {3} The University of Sheffield, United Kingdom

Technical Program – Oct. 27

17:15

785: Apollo: Automated Routing-Informed Placement for Large-Scale Photonic Integrated Circuits

Hongjian Zhou {1}, Haoyu Yang {2}, Gangi Nicholas {3}, Haoxing Ren {4}, Huang Rena {3}, Jiaqi Gu {1}

{1} Arizona State University, United States; {2} NVIDIA Corp., United States; {3} Rensselaer Polytechnic Institute, United States; {4} NVIDIA Corporation, United States

17:30

1231: dyGRASS: Dynamic Spectral Graph Sparsification via Localized Random Walks on GPUs

Yihang Yuan {1}, Ali Aghdaei {2}, Zhuo Feng {1}

{1} Stevens Institute of Technology, United States; {2} University of California San Diego, United States

16:45 - 17:45

Timing Optimization and Optical Routing

Room: Ballroom B

Session Chair(s): Tinghuan CHEN, The Chinese University of Hong Kong, Shenzhen
Wen-Hao Liu, Nvidia

Delivering accuracy and performance in timing, this session begins with studies on concurrent clock buffer sizing and placement refinement, and on buffered Steiner tree construction through reinforcement learning. Physical design and post-layout optimization processes are improved through improved accuracy of parasitic extraction addressed by the third paper. The automation of photonic waveguide routing increases the adoption rate of photonics in integrated circuit technologies is the focus of the last paper.

16:45

223: DiffCCD: Differentiable Concurrent Clock and Data Optimization

Yuhao Ji {1}, Yuntao Lu {2}, Zuodong Zhang {3}, Zizheng Guo {4}, Yibo Lin {4}, Bei Yu {5}

{1} Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} The Chinese University of Hong Kong, China; {3} School of Integrated Circuits, Peking University, China; {4} Peking University, China; {5} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China

17:00

749: Achieving Simultaneous Buffering and Steiner Tree Synthesis via Harmonic Based Reinforcement Learning

Lin Chen, Yuxuan Li, Hu Ding

University of Science and Technology of China, China

Technical Program – Oct. 27

17:15

348: A Complete Modeling Methodology for Full-chip Parasitic Extraction

Yipei Xu {1}, Zhisheng Zeng {2}, Qing He {1}

{1} School of Electronic and Information Engineering, Tongji University, China; {2} Peng Cheng Laboratory, China

17:30

627: Constraints-aware Adaptive Routing with Hybrid Waveguides for Photonic Integrated Circuits

Yuchao Wu {1}, Xiaofei Yu {1}, Xianyi Feng {1}, Yeyu Tong {2}, Yuzhe Ma {1}

{1} The Hong Kong University of Science and Technology (Guangzhou), China; {2} The Hong Kong University of Science and Technology (Guangzhou), China

18:00 – 19:00

Synopsys Sponsor Session

Room: Ballroom A

Session Chair(s): Ron Duncan, Synopsys

19:00 – 20:30

Welcome Reception & SRC Poster Session

Room: Foyer Ballroom

Technical Program – Oct. 28

7:30 – 8:00

Registration

Room: Foyer Ballroom

8:00 – 9:00

Keynote: Driven by AI – Driving Requires More Compute Than Ever

Hans-Jörg Vögel, BMW Group

Room: Ballroom A

Session Chair(s): Deming Chen, University of Illinois

9:00 – 9:30

Coffee Break | Exhibits

Room: Foyer Ballroom

9:30 – 11:00

Panel: Revolution or Hype? Seeking the Limits of Large Models in Hardware Design

Room: Sydney

Moderator: Grace Li Zhang, TU Darmstadt

Panelists: Qiang Xu, CUHK

Leon Stok, IBM

Rolf Drechsler, University of Bremen

Xi Wang, Southeast University

Igor Markov, Synopsys

Recent breakthroughs in Large Language Models (LLMs) and Large Circuit Models (LCMs) have sparked excitement across the electronic design automation (EDA) community, promising significant advances in circuit design and optimization. Yet, skepticism persists: Are these large AI models truly transformative, or are we experiencing a temporary wave of inflated expectations? This panel brings together leading experts from academia and industry to critically examine the practical capabilities, limitations, and future prospects of large AI models in hardware design. Panelists will debate their reliability, scalability, interpretability, and whether these models meaningfully outperform and/or complement traditional EDA methods—offering attendees fresh insights into one of today's most contentious and impactful technology trends.

Technical Program – Oct. 28

9:30 - 11:00

Accelerating Large Language Models (LLMs)

Room: Ballroom B

Session Chair(s): Zheyu Yan, Zhejiang University

Chen Wu, Ningbo Institute of Digital Twin

This session confronts the immense computational and memory challenges of deploying LLMs through diverse hardware-software co-designs. Contributions include heterogeneous architectures like SparCIM to manage sparsity and innovative quantization methods from OA-LAMA and DWA to handle outliers and improve DSP packing. For advanced models, 3D-MoE introduces a 3D-stacked solution for large-scale Mixture-of-Experts, while other works demonstrate how to adapt existing CNN accelerators for LLM inference and even leverage LLMs for SoC design space exploration.

09:30

200: SparCIM: A Heterogeneous CIM-Based Accelerator for Large Language Models with Contextual and Unstructured Bit Sparsity

Xingyu Xu, Yuan Song, Bo Hu, Peng Zheng, Zihan Zou, Xin Si, Bo Liu

Southeast University, China

09:45

272: Discriminate Weight Approximation for Efficient DSP Packing in LLM Accelerators

Jun Li, Yangyijian Liu, Chang-Wei Shi, Mingyang Li, Wu-Jun Li

Nanjing University, China

10:00

383: 3D-MoE: Accelerating Multi-Expert Activated LLMs on 3D In/Near-Memory Computing Architecture via Hybrid Parallelism

Xinyu Qu {1}, Zehua Zhang {2}, Runnan Xu {1}, Yufei Ma {1}

{1} Peking University, China; {2} Tianjin University, China

10:15

757: LLM-Augmented Multi-Modal Fusion for SoC Design Space Exploration

Donger Luo {1}, Tianyi Li {2}, Xinheng Li {3}, Qi Sun {4}, Cheng Zhuo {4}, Bei Yu {5}, Jingyi Yu {2}, Hao Geng {2}

{1} ShanghaiTech University, China; {2} ShanghaiTech University, China; {3} Shanghai Tech, China; {4} Zhejiang University, China; {5} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China

10:30

1440: OA-LAMA: An Outlier-Adaptive LLM Inference Accelerator with Memory-Aligned Mixed-Precision Group Quantization

Huangxu Chen {1}, Yingbo Hao {2}, Xinyu Chen {1}, Yi Zou {2}

{1} The Hong Kong University of Science and Technology (Guangzhou), China; {2} South China University of Technology, China

Technical Program – Oct. 28

10:45

1470: Enabling Decoder-only Language Model Inference on a CNN Accelerator

Seongwoo Choi {1}, Hyunsu Moh {1}, Changjae Yi {1}, Joon Choi {2}, Soonhoi Ha {1}
{1} Seoul National University, Republic of Korea; {2} Chung-Ang University, Republic of Korea

9:30 - 11:00

AI Analog Circuits Design with Learning and Reasoning

Room: Barcelona

Session Chair(s): Bing Li, Universität Siegen

Hailong Yao, University of Science and Technology Beijing

This AI-flavored session showcases the cutting-edge convergence of artificial intelligence and analog circuit design automation. From goal-conditioned reinforcement learning that masters PVT-aware sizing to transformer models achieving zero-shot circuit evaluation, these works represent the forefront of intelligent design methodologies. Witness how knowledge graphs enable topology-agnostic design transfer, reasoning agents revolutionize transistor sizing, and diffusion models generate complete circuit structures. The session further entails robust Bayesian optimization techniques that ensure reliable performance under process variations, demonstrating how AI is fundamentally transforming the analog design landscape.

09:30

250: PPAAS: PVT and Pareto Aware Analog Sizing via Goal-conditioned Reinforcement Learning

Seunggeun Kim {1}, Ziyi Wang {2}, Sungyoung Lee {3}, Youngmin Oh {4}, Hanqing Zhu {1}, Doyun Kim {5}, David Z. Pan {1}

{1} University of Texas at Austin, United States; {2} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {3} The University of Texas at Austin, United States; {4} Samsung Advanced Institute of Technology, Republic of Korea; {5} AI center, Samsung Electronics, Republic of Korea

09:45

412: ZEROSIM: Zero-Shot Analog Circuit Evaluation with Unified Transformer Embeddings

Xiaomeng Yang, Jian Gao, Yanzhi Wang, Xuan Zhang
Northeastern University, United States

10:00

1149: AI Analog Circuit Design Agents: On Knowledge Extraction and Transfer with Knowledge Graphs

Karthik Somayaji N.S., Peng Li
University of California, Santa Barbara, United States

Technical Program – Oct. 28

10:15

1458: ASTRA: Automatic Sizing of Transistors with Reasoning Agents

Wei Xing {1}, Baowen Ou {2}, Yuxuan Zhang {3}, Zhuohua Liu {3}, Yuanqi Hu {3}

{1} The University of Sheffield, United Kingdom; {2} Shenzhen University, China; {3} Beihang University, China

10:30

643: DiffCkt: A Diffusion Model-Based Hybrid Neural Network Framework for Automatic Transistor-Level Generation of Analog Circuits

Chengjie Liu {1}, Jiajia Li {1}, Yabing Feng {1}, Wenhao Huang {2}, Weiyu Chen {3}, Yuan Du {1}, Jun Yang {4}, Li Du {1}

{1} Nanjing University, China; {2} Nanjing University, China; {3} Nanjing University; NCTIEDA, China; {4} Southeast University, China

10:45

1080: RSizing: Robust Bayesian Optimization for Analog Circuit Sizing Under Process Variations

Jindong Tu {1}, Yapeng Li {1}, Peng Xu {2}, Tuo Li {3}, Guoqing Li {3}, Zushuai Xie {4}, Bei Yu {2}, Tinghuan Chen {1}

{1} The Chinese University of Hong Kong, Shenzhen, China; {2} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {3} Shandong Yunhai Guochuang Cloud Computing Equipment Industry Innovation Co., Ltd., China; {4} Nanjing University of Posts and Telecommunications, China

Technical Program – Oct. 28

9:30 - 11:00

Breaking the Thermal Wall: Emerging Thermal Simulation and Design for 2.5D/3D Systems

Room: Calgary

Session Chair(s): Yukai Chen, IMEC, Belgium

Vojtěch Mrázek, Brno University of Technology, Czech Republic

As thermal constraints increasingly define the design envelope for 2.5D/3D and heterogeneously integrated systems, the need for breakthrough thermal solutions has never been greater. This session brings together six pioneering works that fundamentally rethink how we model, simulate, and architect for heat in today's most demanding scenarios, from LLM accelerators to chiplet-based packages. Topics include ultra-fast simulation techniques, geometry-adaptive thermal modeling, and cross-stack thermal-aware architecture. Whether you're designing next-generation compute platforms or optimizing thermal paths across silicon, interposer, and packaging, this session offers actionable insights for researchers and engineers working on thermal-aware system design, chiplet integration, and advanced packaging.

09:30

713: Adaptive Graph Learning for Efficient Thermal Analysis of Multi-Stacking Chiplet Systems under Interface Variations

Ziyao Yang {1}, Yu Cao {1}, Jingbo Sun {2}, Vidya Chhabria {2}

{1} UMN, United States; {2} ASU, United States

09:45

1103: Tasa: Thermal-aware 3D-Stacked Architecture Design with Bandwidth Sharing for LLM Inference

Siyuan He {1}, Peiran Yan {2}, Yandong He {2}, Youwei Zhuo {2}, Tianyu Jia {2}

{1} Beijing, China; {2} Peking University, China

10:00

906: LCTMwalk: GPU-Accelerated Transient Thermal Simulation for Liquid-Cooled 2.5D/3D ICs via Random Walks on Circuit Networks of Modified Compact Thermal Models

Zhixuan Dong {1}, Yonghan Luo {1}, Yuan Meng {2}, Changhao Yan {3}, Zhaori Bi {1}, Keren Zhu {1}, Sheng-Guo Wang {4}, Dian Zhou {5}, Xuan Zeng {1}

{1} Fudan University, China; {2} School of Microelectronics, State Key Laboratory of Integrated Chips & System, Fudan University, China; {3} Associate Prof. Fudan University, China; {4} University of North Carolina at Charlotte, United States; {5} Fudan University, United States

Technical Program – Oct. 28

10:15

22: NSTherm: An Error-Bounded Network-Stochastic Fusion Thermal Simulator for Geometry-Adaptable Chiplets via Diffeomorphic Mapping and Neural-Guided Variance Reduction

Yuan Meng {1}, Yuyan Wang {2}, Zhixuan Dong {3}, Yonghan Luo {3}, Changhao Yan {4}, Keren Zhu {3}, Zhaori Bi {3}, Shengguo Wang {5}, Dian Zhou {3}, Xuan Zeng {3}

{1} School of Microelectronics, State Key Laboratory of Integrated Chips & System, Fudan University, China; {2} State Key Laboratory of Integrated Chips and Systems, Microelectronics Department, Fudan University, China; {3} Fudan University, China; {4} Associate Prof. Fudan University, China; {5} University of North Carolina at Charlotte, United States

10:30

660: High-Resolution Full-Chip Thermal Resistance Extraction of BEOL Interconnects in 3-D ICs Considering Detailed Via Connectivity

Tianxiang Zhu, Qipan Wang, Yibo Lin, Runsheng Wang
Peking University, China

10:45

866: A Geometry-Material Aware Point Cloud Transformer for Large-scale Unstructured Thermal Analysis in 2.5D ICs

Dekang Zhang {1}, Dan Niu {1}, Yichao Cao {1}, Yichao Dong {2}, Zhenya Zhou {3}, Zhou Jin {4}

{1} Southeast University, China; {2} Southeast university, China; {3} Huada Emphyrean Software Co. Ltd, Beijing, China, China; {4} Zhejiang University, China

9:30 - 11:00

Perfecting Routing by All Means

Room: Ballroom A

Session Chair(s): Stephan Held, University of Bonn

Taewhan Kim, Seoul National University

Increased number of components at the PCB, IC, and package levels, as well as accessibility of pins at physical design and post-layout optimization stages, demand sophisticated routing algorithms and strategies. This section includes studies that leverage GPU parallelization and machine learning for acceleration, pin assignment and tuning for accessibility, and routing solutions for 2D IC, 3D IC, and high speed package designs.

09:30

95: GTA: GPU-Accelerated Track Assignment with Lightweight Lookup Table for Conflict Detection

Chunyuan Zhao, Jiarui Wang, Xun Jiang, Jincheng Lou, Yibo Lin
Peking University, China

Technical Program – Oct. 28

09:45

533: H3D: Heterogeneous Resources Aware Global Router for Face-to-Face Bonded 3D ICs

Yuxuan Zhao {1}, Feng Gu {2}, Siting Liu {1}, Peiyu Liao {1}, Bei Yu {1}

{1} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} State Key Lab of Processors, Institute of Computing Technology, Chinese Academy of Sciences; University of Chinese Academy of Sciences; CASTEST, Beijing, China

10:00

371: Adaptive Pin Pattern Modification on Standard Cells Towards ECO Routing

Jaehoon Ahn, Sehyeon Chung, Taewhan Kim

Seoul National University, Republic of Korea

10:15

1155: COPA: A Congestion-Oriented Pin Assignment Framework for Intra-Block Physical Design Optimization

Shu-Yi Tsai {1}, Yu-Guang Chen {1}, Kun-Min Chen {1}, Sheng-Bing Ke {2}, Chung-Hui Hsieh {2}, Chih-Wei Lin {2}, Mango Chia-Tso Chao {3}

{1} National Central University, Taiwan; {2} Realtek Semiconductor Corp., Taiwan; {3} National Yang Ming Chiao Tung University, Taiwan

10:30

598: Performance-Driven Pre-Assignment Routing for High-Speed Package Designs

Shao-Hsiang Chen {1}, Zheng-Wei Chen {1}, Po-Jen Lin {1}, Hung-Jen Hsu {2}, Hsin-Ying Lin {2}, Yung-Hsiang Chuang {1}, Huang-Yu Chen {3}, Jim Chang {2}, Yao-Wen Chang {1}

{1} National Taiwan University, Taiwan; {2} TSMC, Taiwan; {3} TSMC

10:45

971: Refinement Strategies for Any-Angle Package Routing with I/O Alignment Consideration

Yu-En Lin, Shao-Yun Fang, Yi-Yu Liu

National Taiwan University of Science and Technology, Taiwan

Technical Program – Oct. 28

9:30 - 11:00

Special Session: Circuit and Architecture Design with Emerging Computing Paradigms

Room: Atlanta

Session Chair(s): Salim Ullah, Ruhr-Universität Bochum

As emerging computing paradigms push beyond the limitations of traditional CMOS-based computing using Von Neumann architectures, there is a growing need to rethink and extend Electronic Design Automation (EDA) methodologies to support their unique characteristics. These paradigms—including Approximate Computing, In-Memory Computing, Reconfigurable Field-Effect Transistors (RFETs), and Photonic Computing—represent diverse and promising directions beyond conventional digital design. Collectively, they offer transformative potential for achieving significant improvements in energy efficiency, computational speed, and architectural scalability. For example, application-specific approximate computing enables the design of custom arithmetic circuits that exploit application-level error resilience, allowing for optimized accuracy–power–performance–area (PPA) trade-offs in error-tolerant applications. Similarly, the intrinsic polymorphism of RFETs supports compact, multifunctional logic gates and introduces new opportunities for circuit-level obfuscation and security-aware design. Processing-in-non-volatile memories, such as those based on ferroelectric field-effect transistor (FeFETs), enhances energy efficiency by enabling analog computation—particularly for operations like matrix multiplication—directly within the memory arrays. Likewise, photonic analog wavefront computing offers substantial gains in latency and energy efficiency by encoding and processing information in the analog optical domain, leveraging phenomena such as diffraction and interference to perform computation at the speed of light. However, they also introduce a host of new challenges in circuit and architecture design, such as vast and irregular design spaces, analog and non-Boolean behavior, and new device-level constraints that existing EDA tools are not capable of handling. To this end, this special session focuses on developing efficient and robust EDA frameworks that can enable the practical realization of circuits and architectures in these emerging domains.

09:30

Enhancing Robustness of Analog Content Addressable Memory

X. Sharon Hu

University of Notre Dame, United States

09:45

Novel Computing Paradigm Based on Wavefront Analog Information with Photonic-Electronic Hardware for Energy-Efficient Computing

Samarth Vadia

Linque, Germany

Technical Program – Oct. 28

10:00

Secure Hardware Design with Reconfigurable Nano-Structures: An EDA Framework Approach

Akash Kumar

TU Dresden, Germany

10:15

Searching for Application-Specific Approximate Arithmetic Operators: Traditional and ML-Augmented DSE Methods

Siva Satyendra Sahoo

IMEC, Belgium

11:00 – 12:30

Lunch

Room: Ballroom A

12:30 - 14:00

Student Research Competition (Oral Presentations)

Room: Sydney

Session Chair(s): Jiafeng (Harvest) Xie, Villanova University

12:30

5: Beyond Trial and Error: Reliable and Efficient Microfluidics via Systematic Design

Si yuan Liang

The Chinese University of Hong Kong, China

12:39

7: Massively Parallel Logic Synthesis and Verification

Tianji Liu

The Chinese University of Hong Kong, China

12:48

11: Gau++: Algorithm and Hardware Co-Optimized Accelerator Towards Efficient 3D Gaussian Splatting Deployment and Application

Lizhou Wu

Fudan University, China

12:57

12: Design Tools for Adiabatic Superconducting Logic Circuits Toward Energy-Efficient Computing

Rongliang Fu

The Chinese University of Hong Kong, China

Technical Program – Oct. 28

13:06

13: Customizing GNNs with EDA Insights

Ziyi Wang

The Chinese University of Hong Kong, China

13:15

15: Secure and Reliable FeFET-Based Circuit and Architecture Designs for Edge AI

Taixin Li, Xueqing Li

Tsinghua University, China

13:24

18: Long-Horizon Generative AI Acceleration on Edge Hardware

Zizhuo Fu, Meng Li

Peking University, China

13:33

20: Design and Automation of Shuttling in Trapped-Ion Quantum Computers

Daniel Schoenberger

Technical University of Munich, Germany

13:42

25: Optimizing Weight Mapping for Energy-Efficient In-Memory CNN Inference

Johnny Rhe

Sungkyunkwan University, Republic of Korea

13:51

27: Advancing AI for EDA: From Task-Specific Supervised Learning to Circuit Foundation Model

Wenji Fang

Hong Kong University of Science and Technology, China

Technical Program – Oct. 28

12:30 - 14:00

Addressing noise and optimization of quantum and photonic systems

Room: Ballroom A

Session Chair(s): Ajay Joshi, Boston University
Takashi Sato, Kyoto University

Noise and prevention of it is a fundamental challenge in quantum computing. This session will explore new directions in mitigating and correcting errors in quantum systems. Besides that, this session will also include accelerator design, modeling, analysis and optimization of photonic systems.

12:30

417: Revisiting Noise-adaptive Transpilation in Quantum Computing: How Much Impact Does it Have?

Yuqian Huo {1}, Jinbiao Wei {1}, Christopher Kverne {2}, Mayur Akewar {2}, Janki Bhimani {2}, Tirthak Patel {1}

{1} Rice University, United States; {2} Florida International University, United States

12:45

515: SOME: Symmetric One-Hot Matching Elector --- A Lightweight Microsecond Decoder for Quantum Error Correction

Xinyi Guo {1}, Geguang Miao {2}, Shinichi Nishizawa {3}, Hiromitsu Awano {1}, Shinji Kimura {2}, Takashi Sato {1}

{1} Kyoto University, Japan; {2} Waseda University, Japan; {3} Hiroshima University, Japan

13:00

845: EPiCarbon: A Carbon Modeling Tool for Electro-Photonic Accelerators

Farbin Fayza {1}, Cansu Demirkiran {1}, Satyavolu Papa Rao {2}, Darius Bunandar {3}, Udit Gupta {4}, Ajay Joshi {1}

{1} Boston University, United States; {2} NY CREATES, United States; {3} Lighmatter, United States; {4} Cornell Tech, United States

13:15

495: STMC: Small-Tile Multiple-Copy Compilation for Reliable Measurement-Based Quantum Computing

Rongchao Dong, Zewei Mo, Yingheng Li, Aditya Pawar, Jun Yang, Youtao Zhang, Xulong Tang

University of Pittsburgh, United States

13:30

1213: Waferscale Silicon Photonics Systems: A Cost-Benefit Analysis and Optimization

Robert Bao {1}, Frank Cai {1}, Shuangliang Chen {1}, Ajay Joshi {2}, Darius Bunandar {2}, Rakesh Kumar {3}

{1} University of Illinois Urbana-Champaign, United States; {2} Lighmatter, United States; {3} University of Illinois at Urbana-Champaign, United States

Technical Program – Oct. 28

13:45

987: Routing-Aware Placement for Zoned Neutral Atom-based Quantum Computing

Yannick Stade {1}, Wan-Hsuan Lin {2}, Jason Cong {3}, Robert Wille {1}

{1} Technical University of Munich, Germany; {2} University of California, Los Angeles, United States; {3} UCLA, United States

12:30 - 14:00

Brains and Bits: FPGA Acceleration from AI to SAT Solvers

Room: Calgary

Session Chair(s): Chen Zhang, Shanghai Jiao Tong University

Ravikumar V. Chakaravarthy, MIPS

This session explores innovative FPGA-based accelerators targeting a wide range of workloads, from AI model inference and hyperdimensional computing to spiking neural networks and SAT solving. The papers showcase advancements in performance, energy efficiency, and real-time capability for next-generation reconfigurable systems.

12:30

1167: SpecMamba: Accelerating Mamba Inference on FPGA with Speculative Decoding

Linfeng Zhong {1}, Songqiang Xu {1}, Huifeng Wen {1}, Tong Xie {1}, Qingyu Guo {2}, Yuan Wang {1}, Meng Li {3}

{1} Peking University, China; {2} School of Integrated Circuits, Peking University, China; {3} Institute for Artificial Intelligence and School of Integrated Circuits, Peking University, China

12:45

1036: Hyle: An HLS Framework for Hyperdimensional Computing Accelerators on FPGAs

Caio Vieira {1}, Antonio Carlos Schneider Beck {2}

{1} Federal University of Rio Grande do Sul, Brazil; {2} Universidade Federal do Rio Grande do Sul, Brazil

13:00

634: Optimizing Memory Latency and Bandwidth of Spiking Neural Network Accelerators on FPGA via Sparse Hashing

Shadi Matinizadeh, Lakshmi Varshika M, Anup Das

Drexel University, United States

13:15

342: VeriSAT: the Hardware Design of Modern SAT Solver

Yue Tao {1}, Shaowei Cai {2}

{1} Institute of Software, Chinese Academy of Sciences, China; {2} Institute of Software, Chinese Academy of Sciences Country/Region: China (CN), China

Technical Program – Oct. 28

13:30

337: TETRIS: On-Device Trainable Energy-Efficient FPGA Accelerator for Trustworthy and Real-Time Instance Segmentation

Seungil Lee {1}, Juntae Park {2}, Kwanghyun Koo {2}, gilha lee {3}, Sangbeom Jeong {1}, Junoh Park {2}, Hyun Kim {2}

{1} Seoul National University of Science and Technology, Department of Electrical and Information Engineering, Republic of Korea; {2} Seoul National University of Science and Technology, Republic of Korea; {3} SEOUL NATIONAL UNIVERSITY OF SCIENCE AND TECHNOLOGY, Republic of Korea

13:45

29: Hummingbird: A Smaller and Faster Large Language Model Accelerator on Embedded FPGA

Jindong Li {1}, Tenglong Li {1}, Ruiqi Chen {2}, Guobin Shen {1}, Dongcheng Zhao {1}, Qian Zhang {1}, Yi Zeng {1}

{1} Institute of Automation, Chinese Academy of Sciences, China; {2} Vrije Universiteit Brussel, Belgium

12:30 - 14:00

How to Become a Better “Physical Designer”: Novel Tweaks on Mapping, Partitioning and Placement

Room: Barcelona

Session Chair(s): Nimish Agashiwala, Cadence

Yi-Yu Liu, National Taiwan University of Science and Technology

This session contains new techniques to equip modern tools for better placement and partition quality, including GPU leverage (except speed), new models and parallel programming extension, and learning-based methods.

12:30

860: Leveraging GPU for Better Detailed Placement Quality

Chen-Han Lu {1}, Wen-Hao Liu {2}, Haoxing Ren {3}, Ting-Chi Wang {1}

{1} National Tsing Hua University, Taiwan; {2} Nvidia, Taiwan; {3} NVIDIA Corporation, United States

12:45

168: Efficient Analytical Placement Algorithm with Hybrid Fence Region Constraints Using Non-Newtonian Fluid Model

Jai-Ming Lin {1}, Hung-Wei Hsu {1}, Tan HUANG {1}, Chen-Fa Tsai {2}, De-Shiun Fu {2}, Shih-Cheng Huang {2}

{1} Department of Electrical Engineering, National Cheng Kung University, Taiwan; {2} Department of Physical Design Service, Global Unichip Corporation, Tainan, Taiwan, Taiwan

Technical Program – Oct. 28

13:00

552: Ultrafast Density Gradient Accumulation in 3D Analytical Placement with Divergence Theorem

Peiyu Liao, Yuxuan Zhao, Siting Liu, Bei Yu

The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China

13:15

1111: Parallel Non-Integer Multiple-Cell-Height Node Remapping

Zong-Han Wu, Bo-Ying Huang, Yu-Chen Chen, Yao-Wen Chang

National Taiwan University, Taiwan

13:30

1363: HyperEF 2.0: Spectral Hypergraph Coarsening via Krylov Subspace Expansion and Resistance-based Local Clustering

Hamed Sajadinia, Zhuo Feng

Stevens Institute of Technology, United States

13:45

523: Differentiable Timing-Driven FPGA Placement with Smooth Optimization and ML-Based Delay Calibration

Yu-Kang Lin {1}, Zhili Xiong {2}, David Z. Pan {1}

{1} University of Texas at Austin, United States; {2} The University of Texas at Austin, United States

12:30 - 14:00

Reasoning and Debugging in Hardware Design with LLMs

Room: Ballroom B

Session Chair(s): Xinfei Guo, Shanghai Jiao Tong University

Yu Li, Zhejiang University

This session explores recent advances in applying large language models and graph-based learning to key challenges in hardware design, including logic bug detection, RTL verification, and netlist analysis. From automated debugging and reasoning-guided generation to GNN-based circuit matching and layout pattern selection, these works demonstrate how AI can enhance the robustness and scalability of modern hardware development flows. Emphasis is placed on explainability, correctness, and integration into real-world toolchains.

12:30

1027: HLSDebugger: Identification and Correction of Logic Bugs in HLS Code with LLM Solutions

Jing Wang {1}, Shang Liu {1}, Yao Lu {2}, Zhiyao Xie {2}

{1} HKUST, Hong Kong Special Administrative Region of China; {2} Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China

Technical Program – Oct. 28

12:45

1317: Building Reasoning LLMs for Hardware Design Generation via Function-Aligned Differentiated Revision

Weimin Fu {1}, Shijie Li {2}, Yifang Zhao {2}, Kaichen Yang {3}, Xuan Zhang {4}, Yier Jin {2}, Xiaolong Guo {5}

{1} Kansas State University, United States; {2} University of Science and Technology of China, China; {3} Michigan Technological University, United States; {4} Northeastern University, United States; {5} Electrical and Computer Engineering Department, Kansas State University, United States

13:00

913: From Concept to Practice: an Automated LLM-aided UVM Machine for RTL Verification

Junhao Ye {1}, Yuchen Hu {1}, Ke Xu {1}, Dingrong Pan {2}, Qichun Chen {3}, Jie Zhou {1}, Shuai Zhao {4}, Xinwei Fang {5}, Xi Wang {1}, Nan Guan {6}, Zhe Jiang {7}

{1} Southeast University, China; {2} National Center of Technology Innovation for Electronic Design Automation, China; {3} Shenzhen University, China; {4} Sun Yat-sen University, China; {5} University of York, United Kingdom; {6} City University of Hong Kong, Hong Kong Special Administrative Region of China; {7} South East University, China

13:15

191: Target Circuit Matching in Large-Scale Netlists using GNN-Based Region Prediction

Sangwoo Seo {1}, Jimin Seo {2}, Yoonho Lee {1}, Donghyeon Kim {3}, Hyejin Shin {3}, Banghyun Sung {3}, Chanyoung Park {1}

{1} KAIST, Republic of Korea; {2} Kaist Industrial System Engineering, Republic of Korea; {3} SK hynix, Republic of Korea

13:30

566: When Semi-Supervised LVM Meets Frequency-Based Critical Layout Pattern Selection

Liuke Wang {1}, Shenshuo Yao {2}, SHIHAN Wang {3}, Zhen Wang {1}, Zicheng Huang {4}, Jingyi Yu {1}, Hao Geng {1}

{1} ShanghaiTech University, China; {2} Shanghai Tech, China; {3} ShanghaitechUniversity, China; {4} Shanghaitech, China

13:45

127: A Graph-Based Deep Reinforcement Learning Framework for Quantum Circuit Mapping with Look-Ahead Rewards and Biased Exploration

Yueran Zhao {1}, Kaiqi Li {2}, Ziying Guo {2}, Jialin Zhang {2}

{1} School of Computer Science and Technology, University of Chinese Academy of Sciences, Beijing 100049, China, China; {2} Institute of Computing Technology, Chinese Academy of Sciences, China

Technical Program – Oct. 28

12:30 - 14:00

Special Session: Generative AI for At-Scale and Unconventional Analog/RFIC Design and Automation

Room: Atlanta

Session Chair(s): Taiyun Chi, Rice University

Analog/radio-frequency (RF) integrated circuit (IC) design automation has been a longstanding challenge. Recent breakthroughs in generative artificial intelligence (AI) offer transformative opportunities to tackle complex, large-scale analog/RF IC design tasks beyond traditional methods. This special session gathers leading researchers to showcase cutting-edge advancements, from highly efficient inverse design methods to innovative discoveries of unconventional RF passive devices and active analog topologies. These timely developments will set the stage for scalable and extensive use of generative AI in analog/RF IC design, broadly benefiting researchers and engineers in analog/RF IC design, electronic design automation (EDA), and AI/ML communities. Specifically, the first talk will survey recent advances in generative AI-driven automation for analog/RF IC design, from topology generation and sizing to layout and inverse design, and explore challenges in building a fully automated end-to-end framework. The second talk will introduce an agentic AI framework using collaborating LLM agents for sample-efficient and explainable analog circuit sizing. The third talk will present a multi-agent generative synthesis framework that combines graph-based and diffusion models to automate analog/RFIC topology discovery and inverse design. The final talk will showcase an AI-enabled RFIC design flow that integrates reinforcement learning and inverse design for topology, circuit parameters, and passives, validated on mmWave and sub-THz power amplifiers.

12:30

20024: Invited Paper: Generative AI for Analog and RF IC Design: From Spec to Layout

Hyunsu Chae, Seunggeun Kim, Souradip Poddar, Xiaohan Gao, Sensen Li, David Z. Pan
University of Texas at Austin, United States

12:45

20025: Invited Paper: Agentic LLM Workflow for Reasoning-driven Explainable and Sample-Efficient Analog Circuit Sizing

Mohsen Ahmadzadeh, Kaichang Chen, Georges Gielen
KU Leuven, Belgium

13:00

20026: Invited Paper: Multi-Agent Generative Synthesis for Analog/RF Circuits: from Novel Topology Discovery to Efficient Inverse Design

Shikai Wang {1}, Qiufeng Li {1}, Houbo He {2}, Jian Gao {3}, Zining Wang {3}, Yu Sun {4}, Xuan Zhang {3}, Taiyun Chi {2}, Weidong Cao {1}

{1} The George Washington University, United States; {2} Rice University, United States;

{3} Northeastern University, United States; {4} Johns Hopkins University, United States

Technical Program – Oct. 28

13:15

20027: Invited Paper: End-to-end RFIC Topology Synthesis and Design combining Reinforcement learning and Inverse Design

Kaushik Sengupta, Jonathan Zhou, Emir Ali Karahan, Juho Park
Princeton University, United States

14:00 – 14:30

Coffee Break and Exhibits

Room: Foyer Ballroom

14:30 - 16:00

Efficient Verification and Debugging

Room: Calgary

Session Chair(s): Huawei Li, Chinese Academy of Sciences
Lana Josipovic, ETH Zurich

This session presents recent advances in scalable verification techniques for modern hardware systems. It opens with GROOT, a GPU-accelerated framework that combines graph partitioning and edge re-growth to verify extremely large logic circuits efficiently. This is followed by BMCFuzz, which integrates bounded model checking and fuzzing in a hybrid loop to enhance bug detection and coverage. The third paper, Wit-HW, focuses on improving bug localization by generating effective witness test cases beyond the initial bug-triggering to enhance hardware bug localization. The fourth paper introduces a debugging framework that brings source-level tracing and hot-reloading into hardware workflows. The session then turns to safety-critical systems with a graph neural network-based framework aligned with ISO 26262, which uses explainable AI for efficient and interpretable fault analysis. The session concludes with HLSTester, a large language model-driven approach that accelerates the detection of discrepancies in high-level synthesized designs through guided test generation and slicing-based monitoring.

14:30

1379: GROOT: Graph Edge Re-growth and Partitioning for the Verification of Large Designs in Logic Synthesis

Kiran Thorat {1}, Hongwu Peng {1}, Yuebo Luo {2}, Xi Xie {1}, Shaoyi Huang {3}, Amit Hasan {1}, Jiahui Zhao {1}, Yingjie Li {4}, Zhijie Shi {1}, Cunxi Yu {4}, Caiwen Ding {5}

{1} University of Connecticut, United States; {2} University of Minnesota, United States; {3} Stevens Institute of Technology, United States; {4} University of Maryland, College Park, United States; {5} University of Minnesota Twin Cities, United States

Technical Program – Oct. 28

14:45

220: BMCFuzz: Hybrid Verification of Processors by Synergistic Integration of Bound Model Checking and Fuzzing

Shidong Shen {1}, Jinyu Liu {1}, Weizhi Feng {2}, Fu Song {1}, Zhilin Wu {1}

{1} Key Laboratory of System Software (Chinese Academy of Sciences) & State Key Laboratory of Computer Science, Institute of Software, Chinese Academy of Sciences, University of Chinese Academy of Sciences, China; {2} Institute of Software, Chinese Academy of Sciences, China

15:00

439: Wit-HW: Bug Localization in Hardware Design Code via Witness Test Case Generation

Ruiyang Ma {1}, DaiKang Kuang {1}, Ziqian Liu {2}, Jiaxi Zhang {1}, Ping Fan {3}, Guojie Luo {1}

{1} Peking University, China; {2} Renmin University of China, China; {3} DeePoly Technology Inc., China

15:15

1137: Software-Style Hardware Debugging: A Hardware Generation, Simulation and Debugging Framework

Weiran Liu, Shixuan Chen, Chun Yang, Xianhua Liu

Peking University, China

15:30

219: ISO 26262-Aligned Functional Safety Verification framework with Explainable Graph Neural Networks

Yutao Sun {1}, Jiehua Huang {2}, Xiangping Liao {1}, Zhijun Wang {2}, Liping Liang {1}

{1} Beijing University of Posts and Telecommunications, China; {2} BUPT, China

15:45

932: HLSTester: Efficient Testing of Behavioral Discrepancies with LLMs for High-Level Synthesis

Kangwei Xu {1}, Bing Li {2}, Grace Li Zhang {3}, Ulf Schlichtmann {4}

{1} Chair of Electronic Design Automation, Technical University of Munich, Germany; {2} University of Siegen, Germany; {3} TU Darmstadt, Germany; {4} Technical University of Munich, Germany

Technical Program – Oct. 28

14:30 - 16:00

FPGA and Reconfigurable Accelerators for AI Models

Room: Barcelona

Session Chair(s): Yi Sheng, George Mason University
Qi Sun, Zhejiang University

This session demonstrates the power of FPGAs for accelerating a wide range of modern AI models with bespoke hardware solutions. Papers present optimizations for Transformers (SiST, HRAMTran) and prominent alternatives like RWKV (Robin) and Mamba (ToMamba). The versatility of FPGAs is further highlighted by dedicated accelerators for specialized models, including Mixture-of-Experts LLMs (MoE-OPU) and generative Diffusion Transformers (Diff-DiT).

14:30

492: MoE-OPU: An FPGA Overlay Processor Leveraging Expert Parallelism for MoE-based Large Language Models

Shaoqiang Lu {1}, yangbo wei {2}, Junhong Qian {3}, Chen Wu {4}, Xiao Shi {3}, Lei He {5} {1} Shanghai Jiao Tong University, Shanghai, China and Eastern Institute of Technology, Ningbo, China, China; {2} Shanghai Jiao Tong University, China; {3} Southeast University, Nanjin, China, China; {4} Chiplet CAD and Manufacturing Engineering Research Center of Zhejiang Province?Ningbo Institute of Digital Twin?Eastern Institute of Technology , Ningbo, China, China; {5} Eastern Institute of Technology, Ningbo, China and University of California, Los Angeles, United States

14:45

511: Robin: RWKV accelerator using Block Circulant Matrices based on FPGA

Zeyu Li {1}, Shangkun Li {2}, Chuyi DAI {3}, Chaofang MA {1}, Jiawei Liang {4}, Xin Li {5}, Wei Zhang {2}

{1} The Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China; {2} Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China; {3} Fudan University, China; {4} HKUST, Hong Kong Special Administrative Region of China; {5} Duke University, United States

15:00

544: SiST: Token Similarity and Sparsity Aware Optimization for Transformers on FPGA

Yifan Zhang, Hongji Wang, Genhao Zhang, Jianli Chen, Yu Jun, Kun Wang
Fudan University, China

15:15

876: ToMamba: Towards Token-Efficient Mamba Architecture on FPGA

Kejia Shi, Yue Cao, Yuhang Du, Jianli Chen, Jun Yu, Kun Wang
Fudan University, China

Technical Program – Oct. 28

15:30

902: HRAMTran: A Hybrid-RAM Transformer Accelerator With Dynamic Sparsity Floating-Point CIM and Written-Back Transpose Array

Bojun Zhang, Jinkai Wang, Xianan Zhu, Ziyuan Guo, Zhengkun Gu, Kaili Zhang, Kun Zhang, Weisheng Zhao, Yue Zhang
Beihang University, China

15:45

926: Diff-DiT: Temporal Differential Accelerator for Low-bit Diffusion Transformers on FPGA

Shidi Tang {1}, Pengwei Zheng {2}, Ruiqi Chen {3}, Yuxuan Lv {2}, Bruno Silva {3}, Ming Ling {2}

{1} southeast university, China; {2} Southeast University, China; {3} Vrije Universiteit Brussel, Belgium

14:30 - 16:00

Next-Gen Acceleration: Sparsity, Vectors, and Reconfigurable Intelligence

Room: Ballroom B

Session Chair(s): Meng Li, Peking University
Zhenhua Zhu, Tsinghua University

This session explores emerging architectures and optimization techniques that redefine compute efficiency across various workloads. Presentations include novel binary transformer accelerator design for edge inference and accelerators for sparse matrix multiplication—from scan-window strategies to in-DRAM processing—enabling scalable and cost-efficient performance. The session also delves into tensor program optimization for the RISC-V vector extension using probabilistic methods, and introduces a reconfigurable architecture unifying both private and non-private AI inference. Collectively, these works present a forward-looking view into high-performance, flexible, and privacy-aware embedded computing.

14:30

85: COBRA: Algorithm-Architecture Co-optimized Binary Transformer Accelerator for Edge Inference

Ye Qiao, Zhiheng Chen, Yian Wang, Yifan Zhang, Yunzhe Deng, Sitao Huang
University of California, Irvine, United States

14:45

360: ScanNow: A Scan Window-Based Sparse Matrix Multiplication Accelerator Design

Chaofang MA {1}, Lin JIANG {2}, Zeyu LI {1}, Hanwei FAN {3}, Maolin Wang {1}, Jiang XU {4}, Wei Zhang {5}

{1} The Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China; {2} Northeastern University, China; {3} HKUST, Hong Kong Special Administrative Region of China; {4} The Hong Kong University of Science and Technology (Guangzhou), China; {5} Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China

Technical Program – Oct. 28

15:00

859: SPIMA: Scalable and Cost-Efficient Sparse Matrix Multiplication via Processing in DRAM Array

Tairali Assylbekov {1}, Minsang Yu {2}, Jaewoo Park {1}, Mingon Kim {2}, Seungsu Kim {1}, Jongeun Lee {2}

{1} Ulsan National Institute of Science and Technology, Republic of Korea; {2} Ulsan National Institute of Science and Technology (UNIST), Republic of Korea

15:15

225: Tensor Program Optimization for the RISC-V Vector Extension Using Probabilistic Programs

Federico Peccia {1}, Frederik Haxel {1}, Oliver Bringmann {2}

{1} FZI Research Center for Information Technology, Germany; {2} University of Tübingen / FZI, Germany

15:30

267: RTPU: Unifying Non-Private and Private Inference with Reconfigurable Architecture

Fuping Li {1}, Ying Wang {2}, yinghao yang {1}, Jingxuan Li {3}, Yibo Du {4}, Huawei Li {5}, yinhe han {6}, Hang Lu {5}, Xiaowei Li {7}

{1} State Key Lab of Processors, Institute of Computing Technology, China; {2} State Key Laboratory of Computer Architecture, Institute of Computing Technology, Chinese Academy of Sciences, China; {3} University of science and technology of China, China; {4} Institute of Computing Technology, Chinese Academy of Sciences, University of Chinese Academy of Sciences, China; {5} Institute of Computing Technology, Chinese Academy of Sciences, China; {6} Institute of Computing Technology, Chinese Academy of Sciences, China; {7} ICT, Chinese Academy of Sciences, China

14:30 - 16:00

System-level and Architecture Design Optimisation

Room: Ballroom A

Session Chair(s): Cristina Silvano, Politecnico di Milano

Dirk Stroobandt, Universiteit Gent

This session highlights some system-level, architecture and chip design space exploration and optimisation approaches enhanced with AI models and innovative cost and performance models, addressing tasks such as 3D chip design, system and technology co-optimization, multi-technology-node architectures and instruction set customisation for RISC-V processors

14:30

189: Open3DFlow: An Open-Source EDA Platform for 3D Chip Design with AI enhancement

Yifei Zhu {1}, Dawei Feng {1}, Zhenxuan Luan {1}, Zhangxi Tan {2}

{1} Student of TsingHua University, China; {2} Professor of TsingHua University

Technical Program – Oct. 28

14:45

1401: AccelStack: A Cost-Driven Analysis of 3D-Stacked LLM Accelerators

Chen BAI {1}, Xin Fan {1}, Zhenhua Zhu {2}, Wei Zhang {3}, Yuan Xie {4}

{1} Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China; {2} Tsinghua University, China; {3} The Hong Kong University of Science and Technology Country/Region: Hong Kong (HK), Hong Kong Special Administrative Region of China; {4} The Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China

15:00

437: Orthrus: Dual-Loop Automated Framework for System-Technology Co-Optimization

Yi Ren {1}, Baokang Peng {1}, Chenhao Xue {2}, Kairong Guo {1}, Yukun Wang {1}, Guoyao Cheng {3}, Yibo Lin {1}, Lining Zhang {1}, Guangyu Sun {1}

{1} Peking University, China; {2} School of Integrated Circuits, Peking University, China; {3} Peking University, China

15:15

1441: Quantitative Cost Model and Cost Optimization Methods for Multi-Technology-Node Architecture

Qimin Yuan, Kai Huang, Xiaowen Jiang, Dongliang Xiong
Zhejiang University, China

15:30

1006: Seeing Through Designs: Attention-Based Knowledge Transfer for Preference-Guided Microarchitecture Search

Yiyang Zhao {1}, Xuyang Zhao {1}, Zhaori Bi {1}, Ming Zhu {2}, Qiwei Zhan {3}, Keren Zhu {1}, Fan Yang {1}, Changhao Yan {1}, Dian Zhou {4}, Xuan Zeng {1}

{1} Fudan University, China; {2} Anhui University, China; {3} Zhejiang University, China; {4} The University of Texas at Dallas, United States

15:45

362: Automatic Microarchitecture-Aware Custom Instruction Design for RISC-V Processors

Evgenii Rezunov {1}, Niko Zurstraßen {2}, Lennart Reimann {1}, Rainer Leupers {1}

{1} RWTH Aachen University, Germany; {2} RWTH Aachen Institute for Communication Technologies and Embedded Systems, Germany

Technical Program – Oct. 28

14:30 - 16:00

Special Session: Fueling the Future: Achieving Reliable and Generalizable Data Foundations for LLMs in EDA

Room: Sydney

Session Chair(s): Qi Sun, Zhejiang University

Large language models (LLMs) have demonstrated significant potential for enhancing electronic design automation (EDA), particularly in design interpretation and automation workflows. However, practical deployment remains hindered by fundamental data-related obstacles, notably the scarcity of high-quality and standardized datasets. Unlike other domains benefiting from widely available datasets, EDA data are typically fragmented across diverse tools and abstraction layers. Furthermore, synthetic datasets, while widely used, frequently fail to represent the complexities in realistic hardware design, limiting model reliability and generalization. This session brings together leading experts from both academia and industry to address pressing data challenges across the EDA workflow—from RTL to layout, and from design to manufacturing. The goal is to develop robust data generation methods, establish reliable data infrastructures, and curate representative datasets that reflect the complexity of real-world hardware designs. These efforts are critical to building scalable and trustworthy LLM applications that can support next-generation intelligent EDA tools and tackle the growing complexity of modern hardware systems.

14:30

20006: Invited Paper: Diffusion-Model-Enhanced Layout Pattern Generation for Sub 3nm DFM

Guanglei Zhou, Chen-Chia Chang, Junyao Zhang, Jingyu Pan, Yiran Chen
Duke University, United States

14:45

20007: Invited Paper: LLM4Verilog: Building Large-Scale, High-Quality Data Infrastructure for Verilog Code Generation via Community Efforts

Zhongzhi Yu {1}{2}, Chaojian Li {1}, Yongan Zhang {1}, Mingjie Liu {2}, Nathaniel Pinckney {2}, Wenfei Zhou {2}, Rongjian Liang {2}, Haoyu Yang {2}, Haoxing Ren {2}, Yingyan (Celine) Lin {1}
{1} Georgia Institute of Technology, United States; {2} NVIDIA, United States

15:00

20008: Invited Paper: Unitho: A Unified Multi-Task Framework for Computational Lithography

Qian Jin, Yumeng Liu, Yuqi Jiang, Qi Sun, Cheng Zhuo
Zhejiang University, China

15:15

20009: Invited Paper: LLM-Enhanced GPU-Optimized Physical Design at Scale

Yi-Chen Lu, Hao-Hsiang Hsiao, Haoxing Ren
Nvidia, United States

Technical Program – Oct. 28

14:30 - 16:00

Special Session: Reimagining Scientific Computing with Neuromorphic and Analog Systems

Room: Atlanta

Session Chair(s): Anup Das, Drexel University

Antonino Tumeo, Pacific Northwest National Laboratory

Neuromorphic computing, by mimicking the brain behavior, promises to provide highly energy efficient processing of large amounts of data. These systems can have profound impact in Scientific Computing applications, being them sensing applications that require low latency processing or modeling and simulation. In this special session, we plan to discuss the challenges, needs, and opportunities to realize the necessary co-design stack that will make neuromorphic computing effective for scientific problems. We will discuss exemplar scientific computing applications, promising neuromorphic primitives that could enable mapping such broad set of applications, the analog computing aspects underpinning the realization of such neuromorphic primitives, and highlight the necessity of testbed to drive the development of such systems.

14:30

20039: Invited Paper: Designing and Training Neural Networks for Analog In-Sensor Deployment: A Hardware-Aware Analysis

Mark Horton {1}, Haoxuan Shan {1}, James Kiessling {1}, Huanrui Yang {2}, Yiran Chen {1}, Hai “Helen” Li {1}

{1} Duke University, United States; {2} University of Arizona, United States

14:45

20040: Invited Paper: Neuromorphic Architectures for Scientific Computing: a Structural Characterization Case Study

M. L. Varshika {1}{3}, Jonathan Hollenbach {2}, Nicolas Bohm Agostini {3}, Ankur Limaye {3}, Marco Minutoli {3}, Vito Giovanni Castellana {3}, Joseph Manzano {3}, Anup Das {1}, Mitra Taheri {2}{3}, Antonino Tumeo {3}

{1} Drexel University, United States; {2} Johns Hopkins University, United States; {3} Pacific Northwest National Laboratory, United States

15:00

20041: Invited Paper: Synthesizing Analog & Mixed-Signal Floating-Gate enabled Reconfigurable Fabrics using Analog Standard Cells

Jennifer Hasler, Afolabi Ige, Linhao Yang, Pranav Mathews
Georgia Institute of Technology, United States

15:15

20042: Invited Paper: Analyzing the Robustness of Neuromorphic Computing in the presence of Variability in Non-Volatile Memory

Andreia Podasca, Anup Das
Drexel University, United States

Technical Program – Oct. 28

16:00 – 16:15

Transition

16:15 - 17:15

Architecting the Future: AI-Driven Chip Design and HW-Aware Optimization

Room: Barcelona

Session Chair(s): Sung Woo Chung, Korea University

Hussam Amrouch, Technical University of Munich

This session dives into the frontier of AI-guided hardware design in this session, where machine learning meets physical architecture. From chiplet-level layout via LLM agents to sparse Gaussian modeling of DNN accelerators, these papers blend architectural rigor with intelligent automation. Also featured: novel techniques in convolution mechanisms for scalable inference.

16:15

173: Coflex: Enhancing HW-NAS with Sparse Gaussian Processes for Efficient and Scalable DNN Accelerator Design

Yinhui Ma {1}, Zhehui Wang {2}, Tao Luo {2}, Bo Wang {3}, Tomomasa Yamasaki {3}

{1} Singapore University of Technology and Design (SUTD), Singapore; {2} IHP, Singapore;

{3} SUTD, Singapore

16:30

288: QuickFlow: An Efficient Local Search Method to Map Convolutions on Spatial Architectures

Marco Ronzani, Cristina Silvano

Politecnico di Milano, Italy

16:45

545: H2EAL: Hybrid-Bonding Architecture with Hybrid Sparse Attention for Efficient Long-Context LLM Inference

Zizhuo Fu {1}, Xiaotian Guo {2}, Wenxuan Zeng {2}, Shuzhang Zhong {2}, Yadong Zhang {2},

Peiyu Chen {2}, Runsheng Wang {2}, Le Ye {2}, Meng Li {3}

{1} Peking university, China; {2} Peking University, China; {3} Institute for Artificial Intelligence and School of Integrated Circuits, Peking University, China

17:00

1207: MAHL: Multi-Agent LLM-Guided Hierarchical Chiplet Design with Adaptive Debugging

Jinwei Tang {1}, Jiayin Qin {2}, Nuo Xu {3}, Pragnya Nalla {3}, Yu Cao {3}, Yang (Katie) Zhao {2}, Caiwen Ding {1}

{1} University of Minnesota Twin Cities, United States; {2} University of Minnesota, Twin Cities, United States; {3} University of Minnesota, United States

Technical Program – Oct. 28

16:15 - 17:15

Co-Design and Acceleration for Spiking and In-Memory Neural Systems

Room: Ballroom B

Session Chair(s): Grace Li Zhang, TU Darmstadt

Tanja Harbaum, Karlsruhe Institute of Technology

This session explores cutting-edge advancements in spiking neural networks (SNNs) and analog processing-in-memory (PIM), emphasizing algorithm-hardware co-design for enhanced energy efficiency and computational accuracy in neuromorphic and edge computing systems. This first paper introduces a novel 3D hardware architecture for spiking transformers with mixture-of-experts and multi-head attention, achieving improved energy efficiency and latency through dynamic head pruning and 3D integration. The second paper proposes a software-hardware co-design featuring a hardware-friendly spiking direct feedback alignment algorithm and a pipelined RRAM-based IMC accelerator to enable efficient SNN training on edge devices. The third paper presents SPARTA, a spike-aware token skipping framework with a specialized ReRAM-based CIM architecture that co-optimizes structured sparsity in spiking transformers for significant speed and energy improvements. The fourth paper develops a fast, accurate evaluation methodology for analog PIM systems using a novel robustness metric and non-slicing approach, drastically reducing evaluation time while maintaining high fidelity in modeling error impact on neural network accuracy.

16:15

1298: 3D Acceleration for Mixture-of-Experts and Multi-Head Attention Spiking Transformers with Dynamic Head Pruning

Boxun Xu {1}, Junyoung Hwang {2}, Pruek Vanna-iampikul {3}, Yuxuan Yin {1}, Sung Kyu Lim {4}, Peng Li {1}

{1} University of California, Santa Barbara, United States; {2} Georgia Institute of Technology, United States; {3} Burapha University, Thailand; {4} Georgia Tech, United States

16:30

240: When Pipelined In-Memory Accelerators Meet Spiking Direct Feedback Alignment: A Co-Design for Neuromorphic Edge Computing

Haoxiong Ren {1}, Yangu He {2}, Kwunhang Wong {2}, Rui Bao {1}, Ning Lin {2}, Zhongrui Wang {3}, Dashan Shang {1}

{1} Institute of Microelectronics, Chinese Academy of Sciences, China; {2} The University of Hong Kong, Hong Kong Special Administrative Region of China; {3} School of Microelectronics, Southern University of Science and Technology, China

16:45

656: SPARTA: Spike-Aware Token Skipping Co-Optimization with Heterogeneous ReRAM-CIM Architecture for Spiking Transformer Acceleration

Pinfeng Jiang, Letian Wang, Yilong Fang, Yi Wang, Mingde Zhu, Xiangshui Miao, Xingsheng Wang

Huazhong University of Science and Technology, China

Technical Program – Oct. 28

17:00

244: How Do Errors Impact NN Accuracy on Non-Ideal Analog PIM? Fast Evaluation via an Error-Injected Robustness Metric

Lidong Guo {1}, Zhenhua Zhu {1}, Qiushi Lin {1}, Yuan Xie {2}, Huazhong Yang {1}, Wangyang Fu {1}, Yu Wang {1}

{1} Tsinghua University, China; {2} Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China

16:15 - 17:15

Post-CMOS Frontiers: Quantum, Photonics & Neuromorphic Design Innovations

Room: Calgary

Session Chair(s): Mohamed M. Sabry Aly, Nanyang Technological University
Shaocong Wang, University of Notre Dame

This session surveys CAD and circuit advances that exploit beyond-CMOS devices to push performance and energy boundaries. Topics include demand-aware wavelength routing for photonic NoCs, length-matched placement for two-phase RSFQ logic, thin-film LiNbO₃ interposers, programmable photonic logic fabrics, and a learnable printed spiking-neuron architecture. Together, the papers reveal how quantum, photonic, and neuromorphic hardware can be co-designed with modern EDA flows for large PPA benefits. Designers will gain concrete methodologies for integrating these emerging technologies into full VLSI EDA flows.

16:15

561: FAB: Fast and Demand-Aware Bandwidth Allocation Method for Wavelength-Routed Optical Networks-on-Chip

Liaoyuan Cheng, Mengchu Li, Zhidan Zheng, Tsun-Ming Tseng, Ulf Schlichtmann
Technical University of Munich, Germany

16:30

639: J2Place: A Multiphase Clocking-Oriented Length-Matching Placement for Rapid Single-Flux-Quantum Circuits

Rongliang Fu {1}, Minglei Zhou {2}, Huilong Jiang {3}, Junying Huang {4}, Xiaochun Ye {5}, Tsung-Yi Ho {6}

{1} The Chinese University of Hong Kong, China; {2} SKLP, Institute of Computing Technology, CAS; University of Chinese Academy of Sciences, China; {3} State Key Lab of Processors, Institute of Computing Technology, CAS, Beijing, China, China; {4} SKLP, Institute of Computing Technology, CAS, China; {5} State Key Lab of Processors, Institute of Computing Technology, CAS, China; {6} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China

Technical Program – Oct. 28

16:45

1220: VLSI Design and Experimental Demonstration of Photonic Interposers in Thin-Film Lithium Niobate

Georgios Kyriazidis {1}, Aristotelis Tsekouras {2}, John Davis {1}, Vasilis Pavlidis {2}, Gage Hills {1}

{1} Harvard University, United States; {2} Aristotle University of Thessaloniki, Greece

17:00

300: SpikeSynth: Energy-Efficient Adaptive Analog Printed Spiking Neural Networks

Priyanjana Pal {1}, Alexander Studt {1}, Tara Gheshlaghi {1}, Michael Hefenbrock {2}, Michael Beigl {3}, Mehdi Tahoori {1}

{1} Karlsruhe Institute of Technology, Germany; {2} RevoAI GmbH, Germany; {3} Karlsruhe Institute of Technology (KIT), Germany

16:15 - 17:15

Security for AI, AI for Security

Room: Atlanta

Session Chair(s): Lilas Alrahis, Khalifa University

Ericles Sousa, Cadence Design Systems

This session examines the synergistic and thriving relationship between artificial intelligence and hardware security, showcasing both architectures that protect AI systems and novel applications of AI to solve fundamental security challenges.

16:15

1000: "Energon": Unveiling Transformers from GPU Power and Thermal Side-Channels

Arunava Chaudhuri {1}, Shubhi Shukla {2}, Sarani Bhattacharya {2}, Debdeep Mukhopadhyay {3}

{1} Indian Institute of Technology, Kharagpur, India; {2} Indian Institute of Technology Kharagpur, India; {3} Department of Computer Science and Engineering, Indian Institute of Technology Kharagpur, India

16:30

82: Secure Token Pruning Mechanism and Accelerator for Vision Transformer

Qiuran Li {1}, Jingwei Cai {2}, Fanjin Xu {1}, Jingkui Yang {1}, Fei Zhang {1}, Xingyu Zhang {1}, Zhe Li {1}, Jinjin Shao {1}, Yaohua Wang {1}

{1} National University of Defense Technology, China; {2} Tsinghua University, China

16:45

347: Uranus: Ultra-efficient Acceleration Architecture for the Privacy Inference of Graph Neural Networks

Xicheng Xu {1}, Yinghao Yang {2}, Fuyao Liu {1}, Hang Lu {1}, Xiaowei Li {1}

{1} Institute of Computing Technology, Chinese Academy of Sciences, China; {2} State Key Lab of Processors, Institute of Computing Technology, China

Technical Program – Oct. 28

17:00

1160: Unveiling the Mask: Trusted Semiconductor Manufacturing through Wafer-Level Mask-Set Attestation

Suraag Tellakula {1}, Ching-Yi Chang {1}, Matthew Nigh {1}, Christos Vasileiou {1}, John Carulli {2}, Yiorgos Makris {3}

{1} UT Dallas, United States; {2} GlobalFoundries, United States; {3} The University of Texas at Dallas, United States

16:15 - 17:15

Technology-driven Advanced Node Layout Designs

Room: Ballroom A

Session Chair(s): Ting-Chi Wang, National Tsing-Hua University
Marcus Olbrich, Leibniz Universität Hannover

We need technology-aware optimization in order to further obtain PPA benefits from advanced nodes. In this session, we have works focusing on such technology awareness, including 3D, multi-row, and transistor-level.

16:15

1016: ProtoCellLayout: Prototype-Guided Graph Learning for Accurate and Generalizable Cell Layout PPA Estimation

Zhiyuan Luo {1}, Le Zhou {2}, Zenghui Zhang {2}, Jie Zhou {2}, Zhien Li {2}, Huiqing You {3}, Feng Yao {2}, Zhenyu Zhao {2}

{1} National university of defense technology, China; {2} National University of Defense Technology, China; {3} national university of defense technology, China

16:30

1286: Half-Height Double-Row CFET Standard Cells for Area Optimized Placement in A7 CMOS Node

Halil Kükner, Ji-Yung Lin, Sheng Yang, Lynn Verschueren, Jürgen Bömmels, Anita Farokhnejad, Maarten Van de Put, Odysseas Zografos, Naoto Horiguchi, Geert Hellings, Marie Garcia Bardon, Julien Ryckaert

imec, Belgium

16:45

556: DiSPlace: Diffusion-Sharing-Driven Transistor-Level Placement Beyond Standard-Cell Boundaries for DTCO

Keyu Peng {1}, Yinuo Wu {2}, Zhengzhe Zheng {1}, Hao Gu {1}, Ziran Zhu {2}, chao wang {3}, Yang Jun

{1} Southeast University, China; {2} School of Integrated Circuits, Southeast University, China; {3} teacher, China

Technical Program – Oct. 28

17:00

316: Closing the Gap: Advantages of Block-Level over Gate-Level in 3D IC Design for Advanced Nodes

Min Gyu Park {1}, Pruek Vanna-iampikul {2}, Sung Kyu Lim {3}

{1} Georgia Institute of Technology, United States; {2} Burapha University, Thailand; {3} Georgia Tech, United States

16:15 - 17:15

Special Session: Unlocking the benefits of heterogeneous 3D scaling: a major EDA challenge

Room: Sydney

Session Chair(s): Mehdi Tahoori, Karlsruhe Institute of Technology

Chip technology is a foundational enabler across diverse sectors such as automotive, telecommunications, finance, and healthcare. Continued progress in computing performance, efficiency, and functionality has been driven by advances in technology scaling and architecture. However, five major challenges—scaling, memory, power, sustainability, and cost—are now limiting the effectiveness of traditional CMOS design methods at advanced nodes. These limitations demand a fundamental shift in design and technology approaches to meet the increasing performance and efficiency requirements of future applications. This session introduces CMOS 2.0, a new paradigm driven by system technology co-optimization (STCO), where technology and system design are developed in tandem. CMOS 2.0 breaks with the legacy of monolithic scaling by enabling heterogeneous integration through 2.5D and 3D technologies, including chiplet architectures, wafer backside processing, hybrid bonding, and sequential 3D integration. These innovations allow system designers to assign different chip functions—logic, memory, power delivery—into specialized stacked layers, optimizing for performance, energy, and area across a wide variety of workloads. To fully realize the benefits of CMOS 2.0, a corresponding revolution in Electronic Design Automation (EDA) is required. This session will explore the emerging methodologies and tools needed to exploit the expanded design space offered by advanced heterogeneous integration and STCO. It will bring together leading experts to discuss challenges and opportunities in reshaping EDA for CMOS 2.0, offering ICCAD attendees deep insights and practical strategies to drive innovation in the next era of semiconductor design.

16:15

Introducing CMOS 2.0 revolution: challenges and opportunities for EDA

Julien Ryckaert

IMEC, Belgium

16:30

3D-Integration for Heterogeneous Accelerated Computing: a new scale-up hope

Luca Benini

ETH Zurich

Technical Program – Oct. 28

16:45

Tackling Reliability and Yield Challenges of Heterogeneous 3D scaling

Mehdi Tahoori

Karlsruhe Institute of Technology

18:00 – 19:00

Job Fair

Room: Foyer Ballroom

Technical Program – Oct. 29

7:30 – 8:00

Registration

Room: Foyer Ballroom

8:00 – 9:00

Keynote: End-to-end Open Source Platforms in the era of Domain-Specific Design Automation

Luca Benini, Università di Bologna

Room: Ballroom A

Session Chair(s): Ismail S. K. Bustany, AMD

9:00 – 9:30

Coffee Break | Exhibits

Room: Foyer Ballroom

9:30 - 11:00

Advanced Architectures and Emerging Models

Room: Ballroom B

Session Chair(s): Prachi Shukla, AMD

Grace Li Zhang, TU Darmstadt

This session explores the future of AI hardware beyond mainstream accelerators. It features novel paradigms like BITLUME's photonic computing and specialized accelerators for scientific domains like SAFF for GNN-based molecular simulations and Mamba-X for vision-based State Space Models. The session also covers fundamental optimizations, including QUARK's circuit sharing for nonlinear operations, SiGNoR's memory-efficient graph partitioning for GNNs, and SAGE's noise-robust mapping for LLMs on analog hardware.

09:30

46: QUARK: Quantization-Enabled Circuit Sharing for Transformer Acceleration by Exploiting Common Patterns in Nonlinear Operations

Zhixiong Zhao {1}, Haomin Li {2}, Fangxin Liu {3}, Yuncheng Lu {4}, Zongwu Wang {3}, Tao Yang {5}, Haibing Guan {3}, Li Jiang {2}

{1} Nanyang Technological University, China; {2} Shanghai Jiao Tong University, China; {3} Shanghai Jiaotong University, China; {4} Nanyang Technological University, Singapore; {5} Huawei Technologies Co., Ltd, China

09:45

71: BITLUME: Precision-Flexible Photonic Computing for Ultra-Fast and Energy-Efficient DNN Acceleration

Chengpeng Xia {1}, Haibo Zhang {2}, Hao Zhang {1}, Yawen Chen {3}, Amanda Barnard {2}

{1} University of Otago, New Zealand; {2} Australian National University, Australia; {3} University of New South Wales, Australia

Technical Program – Oct. 29

10:00

170: Mamba-X: An End-to-End Vision Mamba Accelerator for Edge Computing Devices

Dongho Yoon, Gungyu Lee, Jaewon Chang, Yunjae Lee, Dongjae Lee, Minsoo Rhu
KAIST, Republic of Korea

10:15

217: SAFF: Scalable Acceleration of GNN-based Machine Learning Force Fields using Tensor-Aware Hardware for Molecular Simulations

Seohye Ha {1}, Yunki Han {1}, Taehwan Kim {1}, Jiwan Kim {1}, Junyoung Park {1}, Gunhee Park {2}, Lee-Sup Kim {1}
{1} KAIST, Republic of Korea; {2} Samsung Electronics, Republic of Korea

10:30

1085: SAGE: Saliency-Aware Grouping for Efficient Mapping of LLMs on Analog Compute-in-Memory

Yayue Hou {1}, Zhenyu Liu {1}, Garrett Gagnon {1}, Hsinyu Tsai {2}, Kaoutar El Maghraoui {2}, Geoffrey Burr {2}, Liu Liu {1}
{1} Rensselaer Polytechnic Institute, United States; {2} IBM, United States

10:45

1444: SiGNOR: Similarity-based Graph Partitioning and Node Reuse for Memory Efficient GNN Acceleration

Seung Eon Hwang, Hyeon Gwon Kim, Dongwoo Lew, Jongsun Park
Korea University, Republic of Korea

9:30 - 11:00

Flow Forward: Frameworks and Optimizations for Reconfigurable Accelerators

Room: Ballroom A

Session Chair(s): Zhengze Jia, Shandong University
Meng Li, Peking University

This session highlights cutting-edge design frameworks and optimization techniques that enable scalable, efficient execution on FPGAs and coarse-grained reconfigurable architectures (CGRAs). The featured works span flexible 3D FPGA architectural exploration, ILP-based synthesis, dynamic vector-dataflow execution, subgraph-based clustering for heterogeneous platforms, spectrum sensing with adaptive context switching, and TinyML with runtime reconfiguration for hybrid inference.

09:30

1239: LaZagna: An Open-Source Framework for Flexible 3D FPGA Architectural Exploration

Ismael Youssef {1}, Hang Yang {1}, Cong "Callie" Hao {2}
{1} Georgia Tech, United States; {2} Georgia Institute of Technology, United States

Technical Program – Oct. 29

09:45

413: ILP-Driven FPGA Multiplier Synthesis: A Scalable Framework for Area-Latency Co-Optimization

Yao Shangshang {1}, Li KunLong {2}, Shen Li {3}

{1} Independent Researcher, China; {2} Fudan University, China; {3} National University of Defense Technology, China

10:00

564: FRESCO: Efficient Subgraph Enumeration for Scalable Clustering in Heterogeneous CGRAs

Louis Coulon {1}, Adham Ragab {2}, Jason Anderson {2}, Mirjana Stojilovic {3}, Paolo Ienne {3}

{1} École Polytechnique Fédérale de Lausanne (EPFL), Switzerland; {2} University of Toronto, Canada; {3} EPFL, Switzerland

10:15

438: DynVec: An End-to-End Framework for Efficient Vector-Dataflow Execution

Jiangnan Li {1}, Xianfeng Cao {2}, Kaixiang Zhu {2}, Wenbo Yin {2}, Lingli Wang {2}

{1} State Key Laboratory of Integrated Chips and System, Fudan University, China; {2} Fudan University, China

10:30

1350: K-PACT: Kernel Planning for Adaptive Context Switching — A Framework for Clustering, Placement, and Prefetching in Spectrum Sensing

Hasan Suluhan {1}, Jiahao Lin {2}, Serhan Gener {3}, Chaitali Chakrabarti {4}, Umit Ogras {5}, Ali Akoglu {3}

{1} The University of Arizona, United States; {2} University of Wisconsin Madison, United States; {3} University of Arizona, United States; {4} Arizona State University, United States; {5} University of Wisconsin - Madison, United States

10:45

966: R2T-Tiny: Runtime-Reconfigurable Throughput-Optimized TinyML for Hybrid Inference Acceleration on FPGA SoCs

Georgios Mentzos {1}, Valentin Frey {1}, Konstantinos Balaskas {2}, Georgios Zervakis {2}, Joerg Henkel {3}

{1} Karlsruhe Institute of Technology, Germany; {2} University of Patras, Greece; {3} KIT, Germany

Technical Program – Oct. 29

9:30 - 11:00

Multimodal and Generative Intelligence for EDA

Room: Barcelona

Session Chair(s): Yu-guang Chen, National Central University
Hao Geng, ShanghaiTech University

This session showcases cutting-edge research at the intersection of multimodal learning and generative AI for electronic design automation. The presented works span cross-modal circuit representation, generative reasoning, and retrieval-augmented LLM frameworks, addressing tasks such as netlist synthesis, RTL decoding, wafer analysis, and documentation QA. Together, these papers reveal the growing potential of LLMs and multimodal models to understand, generate, and optimize hardware designs with unprecedented autonomy and accuracy.

09:30

278: GenEDA: Towards Generative Netlist Functional Reasoning via Cross-Modal Circuit Encoder-Decoder Alignment

Wenji Fang {1}, Jing Wang {2}, Yao Lu {1}, Shang Liu {2}, Zhiyao Xie {1}

{1} Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China; {2} HKUST, Hong Kong Special Administrative Region of China

09:45

392: MMCircuitEval: A Comprehensive Multimodal Circuit-Focused Benchmark for Evaluating LLMs

Chenchen Zhao {1}, Zhengyuan Shi {1}, Xiangyu Wen {2}, Chengjie Liu {3}, Yi Liu {1}, Yunhao Zhou {1}, Yuxiang Zhao {4}, Hefei Feng {5}, Yinan Zhu {6}, Gwok-Waa Wan {6}, Xin Cheng {7}, Weiyu Chen {3}, Yongqi Fu {7}, Chujie Chen {8}, Chenhao Xue {9}, Ying Wang {10}, Yibo Lin {4}, Jun Yang {7}, Ning Xu {7}, Xi Wang {7}, Qiang Xu {1}

{1} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} The Chinese University of Hong Kong, China; {3} Nanjing University, China; {4} Peking University, China; {5} education, China; {6} National Center of Technology Innovation for Electronic Design Automation, China; {7} Southeast University, China; {8} University of Chinese Academy of Sciences, China; {9} School of Integrated Circuits, Peking University, China; {10} State Key Laboratory of Computer Architecture, Institute of Computing Technology, Chinese Academy of Sciences, China

10:00

423: MM-GRADE: A Multi-Modal EDA Tool Documentation QA Framework Leveraging Retrieval Augmented Generation

Yuan Pu {1}, Zhuolun He {1}, Shutong Lin {1}, Jiajun Qin {2}, Xinyun Zhang {1}, Hairuo Han {1}, Haisheng Zheng {3}, Yuqi Jiang {2}, Cheng Zhuo {2}, Qi Sun {2}, David Z. Pan {4}, Bei Yu {1}

{1} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} Zhejiang University, China; {3} Shanghai Artificial Intelligence Laboratory, China; {4} University of Texas at Austin, United States

Technical Program – Oct. 29

10:15

518: FabThink: A Wafer Analysis Multimodal LLM via Chain-of-Thought-Driven Retrieval Augmentation

Yuqi Jiang {1}, Qian Jin {1}, Xudong Lu {1}, Jinyuan Deng {1}, Hao Geng {2}, Hanming Wu {1}, Qi Sun {1}, Cheng Zhuo {1}

{1} Zhejiang University, China; {2} ShanghaiTech University, China

10:30

1233: DecoRTL: A Run-time Decoding Framework for RTL Code Generation with LLMs

Mohammad Akyash, Kimia Zamiri Azar, Hadi Mardani Kamali

University of Central Florida, United States

10:45

1326: CROP: Circuit Retrieval and Optimization with Parameter Guidance using LLMs

Jingyu Pan {1}, Isaac Jacobson {1}, Zheng Zhao {2}, Tung-Chieh Chen {3}, Guanglei Zhou {1}, Chen-Chia Chang {1}, Vineet Rashingkar {4}, Yiran Chen {1}

{1} Duke University, United States; {2} Synopsys Inc, United States; {3} Synopsys, Taiwan; {4} Synopsys, United States

Technical Program – Oct. 29

9:30 - 11:00

Special Session: Cross-Layer Design for Health-Centric AI Systems: Models, Platforms, and Emerging Hardware

Room: Sydney

Session Chair(s): Tinoosh Mohsenin, Johns Hopkins University
Georgios Zervakis, University of Patras

This special session presents a full-stack perspective on AI-driven technologies for healthcare and wellbeing, spanning from hardware innovation and AI model optimization to intelligent orchestration across edge–fog–cloud systems. It opens with the development of Edge Language Models (ELMs), a new class of compressed LLMs designed for real-time interpretation of radiology findings on edge devices, addressing the pressing need for intelligent, low-latency clinical support at the point of care. Enabling such models to function reliably in dynamic, resource-constrained environments requires system-level intelligence—addressed in the second talk through the Mindful AI framework, which dynamically balances accuracy, energy, and latency to ensure robust and personalized operation in pervasive health applications. However, deploying these intelligent systems in real-world, body-integrated scenarios also demands hardware that goes beyond conventional silicon. The third talk addresses this by presenting an end-to-end design framework for flexible healthcare wearables, offering biocompatible, sustainable, and conformable solutions suited for extreme-edge applications such as smart patches. Bringing the stack full circle, the final talk introduces a multi-agent, LLM-driven EDA flow that automates the design of bendable RISC-V processors—empowering non-experts to design custom hardware.

09:30

20010: Invited Paper: BitMedViT: Ternary Quantized Vision Transformer for Medical AI Assistants on the Edge

Mikolaj Walczak, Uttej Kallakuri, Edward Humes, Xiaomin Lin, Tinoosh Mohsenin
Johns Hopkins University, United States

09:45

20011: Invited Paper: Mindful AI for Pervasive Health and Wellbeing (PHW)

Hamidreza Alikhani {1}, Anil Kanduri {2}, Pasi Liljeberg {2}, Amir M. Rahmani {1}, Nikil Dutt {1}
{1} University of California, United States; {2} University of Turku, Finland

10:00

20012: Invited Paper: Feature-to-Classifer Co-Design for Mixed-Signal Smart Flexible Wearables for Healthcare at the Extreme Edge

Maha Shatta {1}, Konstantinos Balaskas {2}, Paula Carolina Lozano Duarte {1}, Georgios Panagopoulos {3}, Mehdi B. Tahoouri {1}, Georgios Zervakis {2}
{1} Karlsruhe Institute of Technology, Germany; {2} University of Patras, Greece; {3} National Technical University of Athens, Greece

Technical Program – Oct. 29

10:15

20013: Invited Paper: Prompt, Fab, Flex: Agentic LLMs for Flexible Electronics Design

Farshad Firouzi {1}, Bahar Farahani {2}, Polykarpos Vergos {3}, Deepesh Sahoo {1}, Nathaniel Bleier {4}, Krishnendu Chakrabarty {1}

{1} Arizona State University, USA; {2} Shahid Beheshti University, Iran; {3} University of Patras, Greece; {4} University of Michigan, USA

Technical Program – Oct. 29

9:30 - 11:00

Special Session: System-Level Design for Chiplets: Architecture Exploration and Resource Management

Room: Atlanta

Session Chair(s): Wanli Chang, Hunan University

As the chip manufacturing node approaches the diameter of silicon atoms, reduction in transistor size becomes unsustainable, faced with severe challenges in yield, power consumption, performance, & cost. Integrating chiplets into 1 package has emerged as a promising solution to further increase the computing capabilities of chips by a few orders of magnitude. It relies on mature manufacturing processes, hence gaining significant advantages in yield protection & agile design without tape-out, exploits heterogeneity of chiplets to support customized design for different application requirements, & possesses ease of use and power savings through advanced packaging. With the interconnection and packaging technologies being heavily invested to build the physical foundation, the focus is moving upwards to the system level, which determines whether the heterogeneous chiplets can cooperate efficiently without any bottlenecks in resource access. It involves architecture exploration & resource management at the chiplet level, which are both new research directions interacting with each other, while the design of individual chiplets will have converged. This session consists of four technical talks from four leading companies of its kind: Intel for chip design & manufacturing, Synopsys developing design automation tools, Visionary Always as a start-up providing systems that manage all types of resources on heterogeneous chiplets, & Bosch representing end users. The first talk from Prashanth Sakthi, Director of System Technology at Intel, will present an architectural and design space exploration methodology tailored for passive interposer-based 2.5D disaggregated systems. The second talk from Tim Kogel, Senior Director for Technical Product Management at Synopsys, will describe AI-driven design space exploration & analysis for multi-die architecture, involving both methodology & tools. The third talk from Wanli Chang, Founder of Visionary Always, will discuss how to manage access to various kinds of resources on heterogeneous chiplets & resolve contention, under different types of workloads that are often dynamic and sometimes have discrepant requirements. The fourth talk from Falk Rehm, Senior Manager & Senior Expert in the field of embedded AI & HW/SW Co-Design at the Bosch Center of Artificial Intelligence, will report the rise of chiplet-based systems in the automotive sector, emphasizing on rapid performance evaluation in the early design stage. The targeted audience include academics & practitioners working in chiplets, especially with a focus on system architecture, design automation tools, advanced packaging, heterogeneous resource management, & simulation, as well as end users.

09:30

20043: Invited Paper: Architectural and Design Space Exploration Strategies for Disaggregated Systems with Passive Die 2.5D Integration

Gauthaman Murali, Mudit Bhargava, Shairfe M. Salahuddin, Zhichao Chen, Archana Pandey, Srivatsa Rangachar Srinivasa, Prerna Budhkar, Ragh Kuttappa, Vinayak Honkote, Prashanth sakthi, Myung-Hee Na, Tanay Karnik
Intel, United States

Technical Program – Oct. 29

09:45

20044: Invited Paper: AI-Driven Multi-Die Architecture Exploration

Tim Kogel

Synopsys, Germany

10:00

20045: Invited Paper: Resource Management on Heterogeneous Chiplets Systems

Wanli Chang^{{1}{2}}, Yili Guo ^{1}, Weijie Wang ^{1}, Yaqi Yao ^{1}, Fuyang Zhao ^{1}, Yinjie Fang ^{1}, Kuan Jiang ^{1}, Liyun Shang ^{1}

^{1} Hunan University, China; ^{2} Visionary Always Technologies, China

10:15

20046: Invited Paper: Automotive Chiplet System: Rapid Performance Evaluation and Optimized AI Inference

Christoph Schorn, Axel Sauer, Marius Fischer, Ingo Feldner, Thomas Schamm, Falk Rehm
Robert Bosch GmbH, Germany

11:00 – 12:30

Lunch

Room: Ballroom A

Technical Program – Oct. 29

12:30 - 14:00

Advancement for Next-generation IC-Design and Computing Paradigm

Room: Ballroom A

Session Chair(s): Takashi Sato, Kyoto University

Vidya Chhabria, Arizona State University

As the semiconductor industry approaches the physical and economic limits of Moore's Law, the impetus for innovation in IC design and computing paradigms has never been greater. This technical session offers a comprehensive platform for presenting and critically evaluating the latest advances, fostering thought leadership, and inspiring transformative progress. This session consists of a broad range of topics. These include analysis of various technological aspects in the 3D-IC space, Electromigration, and Quantum State Preparation.

12:30

433: ThermoPhoton: Fast 3D Thermal Simulation of Photonic Integrated Circuits via Operator Learning

Weilong Guan {1}, Li Huang {1}, Yuxuan Lin {1}, Yuchao Wu {1}, Yeyu Tong {2}, Yuzhe Ma {1}

{1} The Hong Kong University of Science and Technology (Guangzhou), China; {2} The Hong Kong University of Science and Technology (Guangzhou)), China

12:45

751: 3D CoSim: Coupled Operator Learning-Based Co-Simulator for Transferable 3D-IC Analysis

Youran Wu, Shunjie Chang, Jianli Chen, Jun Yu, Kun Wang

Fudan University, China

13:00

605: SCArmor: Layer-Bit Joint Hardening with a Fast Genetic Optimization for Cost-Efficient and High-Reliable SC-DCNN Circuits

Jihe Wang {1}, Yulu Liu {2}, yubin Zhang {2}, Danghui Wang {3}

{1} Computer Science Department, Northwestern Polytechnical University, China; {2} Northwestern Polytechnical University, China; {3} NPU, China

13:15

467: Equivalent Lumped Element Model for Electromigration Considering Thermal Effects

Hengyi Zhu {1}, Wenjie Zhu {1}, Tianshu Hou {2}, Zhigang Ji {3}, Runsheng Wang {4}, Min Tang {1}, Hai-Bao Chen {2}

{1} Shanghai Jiao Tong University, China; {2} Department of Micro/Nano-electronics, Shanghai Jiao Tong University, China; {3} Shanghai Jiaotong University, China; {4} Peking University, China

Technical Program – Oct. 29

13:30

690: A Precision-Steerable Electromigration Solver with Physics-Informed Adaptive Graph Partitioning

Zhaoyuan Liu, Haodong Lu, Jianli Chen, Jun Yu, Kun Wang
Fudan University, China

13:45

1397: Quantum State Preparation Based on LimTDD

Xin Hong {1}, Chenjian Li {1}, Aochu Dai {2}, Sanjiang Li {3}, Shenggang Ying {1}, Mingsheng Ying {4}

{1} Institute of Software, Chinese Academy of Sciences, China; {2} Tsinghua University, China; {3} UTS, Australia; {4} University of Technology Sydney, Australia

12:30 - 14:00

From Prediction to Closure: Intelligent Techniques for Timing Estimation and Optimization

Room: Ballroom B

Session Chair(s): Iris Hui-Ru Jiang, National Taiwan University
Wei W. Xing, The University of Sheffield

Timing closure and optimization are of the key challenges for high-performance IC design. This session will cover the cutting-edge techniques to address this challenge with AI technologies and GPU computing. The first paper in this session presents a GPU-accelerated differentiable physical optimization framework for gate sizing and buffer insertion. The next three papers present the works for accurate timing prediction at the placement stage based on graph transformer, at the pre-synthesis stage considering non-uniform input arrival times, and during the routing stage based on GNN, respectively. Lastly, two papers on the GPU-accelerated incremental STA algorithm and the multi-corner setup/hold time characterization based on active learning are presented.

12:30

206: Differentiable Physical Optimization

Yufan Du, Zizheng Guo, Runsheng Wang, Yibo Lin
Peking University, China

12:45

826: Enhancing Timing Closure via Spatially Embedded Graph Transformer with Low Power/Area Overhead

Joonyoung Seo {1}, Jonghyeon Nam {1}, Howoo Jang {1}, Yunseok Jung {2}, Seokhyeong Kang {1}

{1} Pohang University of Science and Technology, Republic of Korea; {2} POSTECH, Republic of Korea

Technical Program – Oct. 29

13:00

496: NUA-Timer: Pre-Synthesis Timing Prediction Under Non-Uniform Input Arrival Times

Ziyi Wang {1}, Fangzhou Liu {1}, Tsung-Yi Ho {1}, David Z. Pan {2}, Bei Yu {1}

{1} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} University of Texas at Austin, United States

13:15

210: ChronoTE: Crosstalk-Aware Timing Estimation for Routing Optimization via Edge-Enhanced GNNs

Leilei Jin {1}, Rongliang Fu {2}, Zhen Zhuang {1}, Liang Xiao {3}, Fangzhou Liu {1}, Bei Yu {1}, Tsung-Yi Ho {1}

{1} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} The Chinese University of Hong Kong, China; {3} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China

13:30

618: IncreGPUSTA: GPU-Accelerated Incremental Static Timing Analysis for Iterative Design Flows

Haichuan Liu, Zizheng Guo, Runsheng Wang, Yibo Lin

Peking University, China

13:45

249: SetupKit: Efficient Multi-Corner Setup/Hold Time Characterization Using Bias-Enhanced Interpolation and Active Learning

Junzhuo Zhou {1}, Ziwen Wang {2}, HAOXUAN XIA {3}, Yuxin Yan {3}, Chengyu Zhu {4}, Ting-Jung Lin {2}, Wei Xing {5}, Lei He {3}

{1} UCLA, United States; {2} Ningbo Institute of Digital Twin, Eastern Institute of Technology, China; {3} University of California, Los Angeles, United States; {4} BTD Inc., China; {5} The University of Sheffield, United Kingdom

Technical Program – Oct. 29

12:30 - 14:00

Hardware Support for Security: Confidential Computing with Accelerators

Room: Barcelona

Session Chair(s): Johann Knechtel, New York University Abu Dhabi
Joern Stohmann, Cadence Design Systems

This session explores the application and development of hardware acceleration for confidential computing. It brings together two major pillars of the field: hardware-enforced isolation via Trusted Execution Environments (TEEs), and advanced cryptographic techniques like Homomorphic Encryption (HE) and Zero-Knowledge Proofs (ZKPs).

12:30

62: SNO: Securing Network Function Offloading on FPGA-based SmartNICs in Untrusted Clouds

Yunkun Liao {1}, Jingya Wu {2}, Wenyan Lu {3}, Hang Lu {4}, Xiaowei Li {4}, Guihai Yan {3}
{1} SKLP, Institute of Computing Technology, Chinese Academy of Sciences, Beijing, China; University of Chinese Academy of Sciences, Beijing, China; Zhongguancun Laboratory, Beijing, China, China; {2} SKLP, Institute of Computing Technology, Chinese Academy of Sciences, Beijing, China, China; {3} SKLP, Institute of Computing Technology, Chinese Academy of Sciences, Beijing, China; YUSUR Tech Co., Ltd, Beijing, China, China; {4} SKLP, Institute of Computing Technology, Chinese Academy of Sciences, Beijing, China; Zhongguancun Laboratory, Beijing, China, China

12:45

1313: FPGA-CC: Confidential Containers for Virtualized FPGAs

Ke Xia, Sheng Wei
Rutgers University, United States

13:00

1249: FlexTEE: Dynamically Enhancing Metadata Locality Through Affine Address Transformation for Heterogeneous & Secure AI Platforms

Rakin Muhammad Shadab, Sanjay Gandham, Mingjie Lin
University of Central Florida, United States

13:15

728: PATHE: A Privacy-Preserving Database Pattern Search Platform with Homomorphic Encryption

Xuan Wang {1}, Minxuan Zhou {2}, Gabrielle De Micheli {1}, Yujin Nam {1}, Sumukh Pingre {3}, Augusto Vega {4}, Tajana Rosing {5}
{1} University of California San Diego, United States; {2} Illinois Tech, United States; {3} University of California, San Diego, United States; {4} IBM Research, United States; {5} UCSD, United States

Technical Program – Oct. 29

13:30

867: Making the Best Switch: Encoding Strategy Management for Efficient TFHE Circuit Evaluation

Mingfei Yu {1}, Gabrielle De Micheli {2}, Giovanni De Micheli {1}
{1} EPFL, Switzerland; {2} UCSD, United States

13:45

1290: Gotta Hash 'Em All! Accelerating Hash Functions for Zero-Knowledge Proof Applications

Nojan Sheybani {1}, Teng kai Gong {2}, Anees Ahmed {3}, Nges Brian Njungle {4}, Michel Kinsy {3}, Farinaz Koushanfar {2}
{1} UC San Diego, United States; {2} University of California San Diego, United States; {3} Arizona State University, United States; {4} Arizona State university, United States

12:30 - 14:00

Scalable Simulation and Verification Frameworks

Room: Atlanta

Session Chair(s): Keren Zhu, Fudan University
Kaveh Razavi, ETH Zürich

This session presents a comprehensive collection of techniques for simulation and verification. It begins with an augmented co-simulation framework that accelerates system-level verification of HLS accelerators. The next paper introduces CTDM, a time-division multiplexing method for resource-efficient, FPGA-accelerated simulation of large-scale neural processing unit designs. WROXIM follows with a simulator for wavelength-routed optical NoCs, enabling accurate modeling of communication across many-core systems. COTIA then demonstrates how an intelligent agent-guided concolic testing approach improves test path exploration. Promise presents a property mining strategy that enhances sequential synthesis by eliminating redundant logic based on mined invariants. The session concludes with Pathfinder, which builds cycle-accurate taint graphs to trace and analyze information flow.

12:30

853: Augmented Co-Simulation for Fast Functional and System-Level Verification of HLS Accelerators

Michele Fiorito, Serena Curzel, Fabrizio Ferrandi
Politecnico di Milano, Italy

12:45

727: CTDM: Resource-Efficient FPGA-Accelerated Simulation of Large-Scale NPU Designs

Hyunje Jo {1}, Han-sok Suh {2}, Hyungseok Heo {1}, Jinseok Kim {1}, Hyunsung Kim {1}, Boeui Hong {1}, Jungju Oh {1}, Sunghyun Park {1}, Jinwook Oh {1}, Sunghwan Jo {3}, Kangwook Lee {3}, Jae-sun Seo {4}
{1} Rebellions, Republic of Korea; {2} Cornell University, United States; {3} Synopsys Korea, Republic of Korea; {4} Cornell Tech, United States

Technical Program – Oct. 29

13:00

493: WROXIM: A Network-Level Simulation Platform for Wavelength-Routed Optical Networks-on-Chip

Jeng-De Chang {1}, Zhidan Zheng {2}, Liaoyuan Cheng {2}, Liu-Xuan-Wei Zhang {1}, Tsun-Ming Tseng {2}, Ing-Chao Lin {3}, Ulf Schlichtmann {2}

{1} National Cheng Kung University, Taiwan; {2} Technical University of Munich, Germany; {3} National Yang Ming Chiao Tung University

13:15

213: COTIA: Concolic Testing with Intelligent Agent

Yan TAN {1}, Xiangchen Meng {2}, Yangdi Lyu {3}

{1} HKUST-GZ, China; {2} Hong Kong University of Science and Technology(Guangzhou), China; {3} Hong Kong University of Science and Technology (Guangzhou), China

13:30

998: Promise: Property Mining for Sequential Synthesis

Jiahui Xu {1}, Jordi Cortadella {2}, Lana Josipovic {1}

{1} ETH Zurich, Switzerland; {2} Universitat Politecnica de Catalunya, Spain

13:45

537: Pathfinder: Constructing Cycle-accurate Taint Graphs for Analyzing Information Flow Traces

Katharina Ceesay-Seitz {1}, Flavien Solt {2}, Alexander Klukas {1}, Kaveh Razavi {1}

{1} ETH Zurich, Switzerland; {2} UC Berkeley, United States

12:30 - 14:00

Special Session: 2025 CAD Contests at ICCAD

Room: Sydney

Session Chair(s): Shao-Yun Fang, National Taiwan University of Science and Technology

The CAD Contest at ICCAD (<https://iccad-contest.org/>) is a challenging, multi-month, research & development competition, focusing on advanced, real-world problems in the field of Electronic Design Automation (EDA). Contestants can participate in one or more problems provided by EDA/IC industry. The winners will be awarded at an ICCAD special session dedicated to this contest. Since 2012, the CAD Contest at ICCAD has been attracting more than a hundred teams per year, fostering productive industry-academia collaborations, and leading to hundreds of publications in top-tier conferences and journals. The contest keeps enhancing its impact and boosts EDA research.

12:30

20018: Invited Paper: Overview of 2025 CAD Contest at ICCAD

Chung-Kuan Cheng {1}, Shao-Yun Fang {2}, Yi-Yu Liu {2}, Tsun-Ming Tseng {3}

{1} University of California San Diego, USA; {2} National Taiwan University of Science and Technology, Taiwan; {3} Technical University of Munich, Germany

Technical Program – Oct. 29

12:45

20019: Invited Paper: 2025 ICCAD CAD Contest Problem A: Hardware Trojan Detection on Gate Level Netlist

Chung-Han Chou {1}, Kai-Chiang Wu {2}, Chih-Jen (Jacky) Hsu {3}, Yu-Guang Chen {4}, Hung-Chun Chiu {1}, Zhuo Li {3}

{1} Cadence, Taiwan; {2} National Yang Ming Chiao Tung University, Taiwan; {3} Cadence, United States; {4} National Central University, Taiwan

13:00

20020: Invited Paper: 2025 ICCAD CAD Contest Problem B: Power and Timing Optimization Using Multibit Flip-Flop

Sheng-Wei Yang, Jhih-Wei Hsu, Yu-Hsuan Cheng, Chin-Fang Cindy Shen

Synopsys, Taiwan

13:15

20021: Invited Paper: 2025 ICCAD CAD Contest Problem C: Incremental Placement Optimization Beyond Detailed Placement: Simultaneous Gate Sizing, Buffering, and Cell Relocation

Yi-Chen Lu, Rongjian Liang, Wen-Hao Liu, Haoxing Ren

Nvidia, United States

13:30

20022: Invited Paper: IEEE DATC RDF-2025: Enabling an EDA Research Ecosystem

Vidya A. Chhabria {1}, Amur Ghose {2}, Vikram Gopalakrishnan {1}, Andrew B. Kahng {2}, Sayak Kundu {2}, Yiting Liu {2}, Zhiang Wang {2}, Bing-Yue Wu {1}

{1} Arizona State University, United States; {2} University of California San Diego, United States

14:00 – 14:30

Coffee Break and Exhibits

Room: Foyer Ballroom

Technical Program – Oct. 29

14:30 - 16:00

Next-Gen Application Acceleration Techniques

Room: Atlanta

Session Chair(s): Wei Zhang, The Hong Kong University of Science and Technology
Christian Hochberger, TU Darmstadt

This session showcases some innovative approaches for accelerating OS kernels, such as scheduling and I/O middleware, as well as different applications such as e-graph extraction, genome analysis, and AI model compression.

14:30

803: HERCULES: Hardware accElerator foR stoChastic schedULing for hEterogeneous Systems

Vairavan Palaniappan {1}, Adam Ross {1}, Amit Trivedi {2}, Debjit Pal {2}

{1} University of Illinois Chicago, United States; {2} University of Illinois at Chicago, United States

14:45

1144: GAIA: Glass-Aware I/O Middleware

Hung-Yuan Cheng {1}, Chun-Feng Wu {2}, Yuan-Hao Chang {3}, David Hung-Chang Du {4}

{1} 887000000000, Taiwan; {2} National Yang Ming Chiao Tung University, Taiwan; {3} Academia Sinica, Taiwan; {4} University of Minnesota, United States

15:00

1318: e-boost: Boosted E-Graph Extraction with Adaptive Heuristics and Exact Solving

Jiaqi Yin {1}, Zhan Song {1}, Chen Chen {1}, Yaohui Cai {2}, Zhiru Zhang {2}, Cunxi Yu {1}

{1} University of Maryland, College Park, United States; {2} Cornell University, United States

15:15

425: GenomeDPU: A Cost-Effective In-Memory Data Processing Unit for GPU-based Genome Analysis

Shouzhe Zhang, Zhuren Liu, Ruixiao Huang, Hui Zhao

University of North Texas, United States

15:30

534: CoXplorer: Multi-staged Co-exploration Framework for AI Model Compression and Accelerator Design

Songchen Ma {1}, Junyi Wu {2}, Yonghao Tan {3}, Pingcheng Dong {1}, Peng Luo {4}, Di Pang {4}, Yu Liu {4}, Xuejiao Liu {4}, Luhong Liang {4}, Tim Cheng {5}, Fengbin Tu {1}

{1} The Hong Kong University of Science and Technology, China; {2} School of Electronic science and engineering, Nanjing Univerisity, China; {3} The Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China; {4} ACCESS -- AI Chip Center for Emerging Smart Systems, Hong Kong Special Administrative Region of China; {5} HKUST, Hong Kong Special Administrative Region of China

Technical Program – Oct. 29

14:30 - 16:00

Simulation and Synthesis of quantum circuits and systems

Room: Ballroom A

Session Chair(s): Robert Wille, TU Munich

Vidya Chhabria, Arizona State University

The rapidly growing capacity of Quantum computers need to move in tandem with the design automation tools for the same. This session will cover quantum circuit simulation and synthesis, including consideration for reliability. Techniques like quantization, qudit gates and parallelotopes will be explored. Furthermore, this session will cover the new findings in quantum circuit simulation, synthesis and routing, including various quantum technologies such as superconducting and spin qubits.

14:30

255: QQ: Is 2-bit Enough? Exploiting Quantization to Enhance Computation and Memory Efficiency in Quantum Simulation

Hyoju Seo, Seokhyeon Lee, Yongtae Kim

Kyungpook National University, Republic of Korea

14:45

1094: Simulation Framework and Optimization of Superconducting Transmon-Tunable Coupler-Transmon System for Qudit Gates

Ferris Prima Nugraha, Yuhan Huang, Jiacheng Liu, Qiming Shao

The Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China

15:00

679: CNOT Oriented Synthesis for Small-Scale Boolean Functions Using Spatial Structures of Parallelotopes

Qiang Zheng {1}, Yongzhen Xu {2}, Jiayi Zhang {3}, Zhaofeng Su {1}, Shenggen Zheng {4}

{1} University of Science and Technology of China, China; {2} Quantum Science Center of Guangdong-Hong Kong-Macao Greater Bay, China; {3} Peking University, China; {4} Quantum Science Center of Guangdong-Hong Kong-Macao Greater Bay Area, China

15:15

265: CLASS: A Controller-Centric Layout Synthesizer for Dynamic Quantum Circuits

Yu Chen {1}, Yilun Zhao {2}, Bing Li {3}, He Li {4}, Mengdi Wang {1}, Yinhe Han {5}, Ying Wang {6}

{1} Institute of Computing Technology, Chinese Academy of Sciences, China; {2} Institute of Computing Technology, CAS, China; {3} Capital Normal University, China; {4} Southeast University, China; {5} Institute of Computing Technology, Chinese Academy of Sciences, China; {6} State Key Laboratory of Computer Architecture, Institute of Computing Technology, Chinese Academy of Sciences, China

Technical Program – Oct. 29

15:30

692: Circuit Folding: Scalable and Graph-Based Circuit Cutting via Modular Structure Exploitation

Shuwen Kan {1}, Yanni Li {1}, Hao Wang {2}, Sara Mouradian {3}, Ying Mao {1}
{1} Fordham University, United States; {2} Stevens Institute of Technology, United States;
{3} University of Washington, United States

15:45

632: An Efficient Routing Optimization Framework for Silicon-Based Spin-Qubit Devices

Ching-Yao Huang, Wai-Kei Mak
National Tsing Hua University, Taiwan

14:30 - 16:00

Smarter RTL and Architecture: With or Without AI

Room: Barcelona

Session Chair(s): Christian Pilato, Politecnico di Milano
Cunxi Yu, cunxiyu@umd.edu

From LLM-powered RTL synthesis to architecture-aware exploration and circuit-level learning, this session presents diverse approaches to improving design quality. Whether driven by AI or domain-specific algorithms, these works target better RTL, better architectures, and better insights into digital hardware.

14:30

110: VERIRL: Boosting the LLM-based Verilog Code Generation via Reinforcement Learning

Fu Teng {1}, Miao Pan {1}, Xuhong Zhang {1}, Zhezhi He {2}, Yiyao Yang {2}, Xinyi Chai {1}, Mengnan Qi {2}, Liqiang Lu {1}, Jianwei Yin {1}
{1} Zhejiang university, China; {2} Shanghai Jiao Tong University, China

14:45

1222: VeriOpt: PPA-Aware High-Quality Verilog Generation via Multi-Role LLMs

Kimia Tasnia {1}, Alexander Garcia {1}, Tasnuva Farheen {2}, Sazadur Rahman {1}
{1} University of Central Florida, United States; {2} Louisiana State University, United States

15:00

109: Automated Design Space Exploration in High-Level Physical Synthesis

Linfeng Du {1}, Jiawei Liang {2}, Jason Lau {3}, Yuze Chi {3}, Yutong Xie {4}, Chunyou SU {1}, Afzal Ahmad {5}, Zifan He {6}, Jake Ke {6}, Jinming Ge {5}, Jason Cong {3}, Wei Zhang {5}, Licheng Guo {3}
{1} The Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China; {2} HKUST, Hong Kong Special Administrative Region of China; {3} UCLA, United States; {4} RapidStream Design Automation, Inc., United States; {5} Hong Kong University of Science and Technology, Hong Kong Special Administrative Region of China; {6} University of California, Los Angeles, United States

Technical Program – Oct. 29

15:15

468: Clay: High-level ASIP Framework for Flexible Microarchitecture-Aware Instruction Customization

Weijie Peng {1}, Youwei Xiao {1}, Yuyang Zou {2}, Zizhang Luo {1}, Yun (Eric) Liang {2}
{1} School of Integrated Circuits, Peking University, China; {2} Peking University, China

15:30

415: Optimizing SFQ Circuit Design: A Timing-Driven Framework for Performance-Constrained Area Minimization

Robert Aviles {1}, Rassul Bairamkulov {2}, Ziyu Liu {1}, Peter Beerel {1}
{1} University of Southern California, United States; {2} Advanced Micro Devices Inc., United States

15:45

825: DeepCell: Self-Supervised Multiview Fusion for Circuit Representation Learning

Zhengyuan Shi {1}, Chengyu Ma {2}, Ziyang Zheng {1}, Lingfeng Zhou {3}, Hongyang Pan {4}, Wentao Jiang {2}, Fan Yang {5}, Xiaoyan Yang {3}, Zhufei Chu {2}, Qiang Xu {1}
{1} The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China; {2} Ningbo University, China; {3} Hangzhou Dianzi University, China; {4} Fudan university, China; {5} Fudan University, China

14:30 - 16:00

Solving scientific problems with chiplets, Flash memory and elastic computing

Room: Ballroom B

Session Chair(s): Martin Margala, University of Louisiana at Lafayette
Gage Hills, Harvard University

Rapid evolution of application-driven demands are nowadays quickly connected with technological innovations at device/circuit level. This session will present the solutions to pressing scientific problems of current times such as sustainability, cybersecurity, search and genome sequencing by exploiting bleeding edge technologies. Examples of this include elastic computing to efficiently compute complex arithmetic primitives, and accelerating LLM inference using 3D chiplet. Besides these cross-stack works, sustainability issues, and chiplet-specific physical design problems will also be discussed.

14:30

50: RePM: Reconfigurable Elastic Computing for Polynomial Multiplier with Hybrid NTT Algorithm

Sizhao Li {1}, Chenyu Zhai {1}, Xinhua Wang {1}, Zhujun Guo {1}, Shan He {2}, Donghui GUO {2}
{1} Harbin Engineering University, China; {2} Xiamen University, China

Technical Program – Oct. 29

14:45

405: FeNOMS: Enhancing Open Modification Spectral Library Search with In-Storage Processing of Ferroelectric NAND (FeNAND) Flash

Sumukh Pinge {1}, Ashkan Moradifirouzabadi {1}, Keming Fan {1}, Prasanna Venkatesan Ravindran {2}, Tanvir Pantha {3}, Po-Kai Hsu {4}, Zheyu Li {5}, Weihong Xu {6}, Zihan Xia {7}, Flavio Ponzina {7}, Winston Chern {2}, Taeyoung Song {2}, Priyanka Ravikumar {2}, Mengkun Tian {2}, Lance Fernandes {2}, Huy Tran {2}, Hari Jayasankar {2}, Hang Chen {2}, Chinsung Park {2}, Amrit Garlapati {2}, Kijoon Kim {8}, Jongho Woo {8}, Suhwan Lim {8}, Kwangsoo Kim {8}, Wanki Kim {8}, Daewon Ha {8}, Duygu Kuzum {7}

{1} University of California, San Diego, United States; {2} Georgia Institute of Technology, United States; {3} The University of Texas at Dallas, United States; {4} Macronix, Taiwan; {5} UCSD, United States; {6} EPFL, Switzerland; {7} University of California San Diego, United States; {8} Semiconductor Research and Development, Samsung Electronics Co., Ltd, Republic of Korea

15:00

334: H3D-LLM: Heterogeneous 3D Chiplet Design for LLM Inference with Dynamic Task Scheduling and Memory-Aware Orchestration

Hui Kou, Chenjie Xia, Jialin Yang, Liyi Li, Hao Cai, Xin Si, Bo Liu
Southeast University, China

15:15

416: CarbonClarity: Understanding and Addressing Uncertainty in Embodied Carbon for Sustainable Computing

Xuesi Chen {1}, Leo Han {2}, Anvita Bhagavathula {1}, Udit Gupta {1}
{1} Cornell Tech, United States; {2} Cornell University, United States

15:30

810: Accelerating Genome Alignment Pipeline with In-NAND Search Technology and Group Testing Techniques

Ming-Hsiang Tsai {1}, Ming-Liang Wei {2}, Chia-Chun Chien {1}, Po-Hao Tseng {3}, Yung-Chun Lee {4}, Hsiang-Pang Li {4}, Chia-Lin Yang {2}

{1} Graduate Institute of Electronics Engineering, National Taiwan University, Taiwan; {2} National Taiwan University, Taiwan; {3} Macronix International co., Ltd., Taiwan; {4} Macronix International co., Ltd., Taiwan

15:45

330: P2P-Chiplet: Partition and Placement Co-Optimization for Multi-Chiplet Architecture

Qidie Wu {1}, Jiangyuan Gu {2}, Xuguang Yuan {1}, Shaojun Wei {2}, Shouyi YIN {2}

{1} School of Integrated Circuits, Tsinghua University, China; {2} Tsinghua University, China

Technical Program – Oct. 29

14:30 - 16:00

Special Session: Privacy-Preserving Deep Learning in the LLM Era: Algorithm, Protocol, and Hardware

Room: Sydney

Session Chair(s): Meng Li, Peking University
Wujie Wen, North Carolina State University

The last two years have witnessed the rapid evolution of large language models (LLMs). While data serves as the foundation for such evolution, it also faces ever-increasing privacy risks, including data breaches, leakage, misuse, etc, as existing LLM services on the cloud heavily rely on the availability of plaintext prompts. Privacy-preserving deep learning (PPDL) has thus been proposed and has recently attracted increasing attention from both industry and academia. By leveraging cryptographic primitives such as homomorphic encryption (HE), multi-party computation (MPC), etc, PPDL achieves a formal privacy guarantee during deep learning (DL) computation. However, PPDL usually suffers from orders of magnitude latency overhead due to high computation and communication costs, hindering its practical usage in real-world applications. This special session presents four in-depth talks that cover a full stack of techniques across the algorithm, protocol, and hardware levels to narrow the efficiency gap of PPDL. By fostering an interdisciplinary dialogue, this session not only presents to the audience the recent progress of PPDL but also encourages discussion on the challenges and opportunities of PPDL.

14:30

20014: Invited Paper: Optimizing Privacy-Preserving Primitives to Support LLM-Scale Applications

Yaman Jandali, Ruisi Zhang, Nojan Sheybani, Farinaz Koushanfar
University of California, San Diego, United States

14:45

20015: Invited Paper: FENIX: Flexible and Efficient Hybrid HE/MPC Acceleration with Near-Memory Processing

Tengyu Zhang {1}, Chenqi Lin {1}, Jiangrui Yu {1}, Yi Chen {1}, Shuwen Deng {2}, Meng Li {1}{3}
{1} Peking University, China; {2} Tsinghua University, China; {3} Beijing Advanced Innovation Center for Integrated Circuits, China

15:00

20016: Invited Paper: Network and Compiler Optimizations for Efficient Linear Algebra Kernels in Private Inference

Karthik Garimella, Negar Neda, Austin Ebel, Nandan Kumar Jha, Brandon Reagen
New York University, United States

Technical Program – Oct. 29

15:15

20017: Invited Paper: Maximizing SIMD Efficiency for Encrypted Machine Learning via Pattern-Aware Encoding

Ran Ran, Zhaoting Gong, Zhaowei Li, Wujie Wen
North Carolina State University, United States

16:00 – 16:15

Transition

16:15 - 17:15

Efficient AI on Edge and LLM Infrastructure

Room: Ballroom A

Session Chair(s): Jason Xue, MBZUAI
Albert Zeng, Cadence

This session highlights innovations in deploying efficient AI models on edge devices and improving the scalability of LLM infrastructure. The selected works address challenges in model sparsity, quantization, cross-modal alignment, and near-sensor accelerator with silicon photonics. By bridging high-performance learning with energy-efficient execution, these papers pave the way for the next generation of compact, responsive AI systems tailored to resource-constrained environments.

16:15

933: LLM-Barber: Block-Aware Rebuilder for Sparsity Mask in One-Shot for Large Language Models

Yupeng Su {1}, Ziyi Guan {2}, Xiaoqun Liu {1}, Tianlai Jin {1}, Dongkuan Wu {1}, Zhengfei Chen {1}, Graziano Chesi {3}, Ngai Wong {2}, Hao Yu {1}
{1} Southern University of Science and Technology, China; {2} The University of Hong Kong, Hong Kong Special Administrative Region of China; {3} The University of Hong Kong, Italy

16:30

75: Tiny-Align: Bridging Automatic Speech Recognition and Large Language Model on Edge

Ruiyang Qin {1}, Dancheng Liu {2}, Gelei Xu {3}, Amir Nassereldine {2}, Zheyu Yan {4}, Chenhui Xu {2}, yuting hu {5}, Shaocong Wang {3}, X. Sharon Hu {3}, Jinjun Xiong {2}, Yiyu Shi {3}
{1} Villanova University, United States; {2} University at Buffalo, United States; {3} University of Notre Dame, United States; {4} Zhejiang University, China; {5} The State University of New York at Buffalo, United States

Technical Program – Oct. 29

16:45

105: Squat: Quant Small Language Models on the Edge

Xuan Shen {1}, Peiyan Dong {2}, Zhenglun Kong {1}, Yifan Gong {1}, Changdi Yang {1}, Zhaoyang Han {1}, Yanyue Xie {1}, LEI LU {1}, Cheng Lyu {3}, Chao Wu {1}, Yanzhi Wang {1}, PU ZHAO {4}

{1} Northeastern University, United States; {2} Massachusetts Institute of Technology (MIT), United States; {3} AlBao, United States; {4} Northeastern.edu, United States

17:00

1309: Opto-ViT: Architecting a Near-Sensor Region of Interest-Aware Vision Transformer Accelerator with Silicon Photonics

Mehrdad Morsali {1}, Chengwei Zhou {2}, Deniz Najafi {3}, Sreetama Sarkar {4}, Pietro Mercati {5}, Navid Khoshavi {6}, Peter Beerel {7}, Mahdi Nikdast {8}, Gourav Datta {2}, Shaahin Angizi {3}

{1} New jersey Institute of Technology, United States; {2} Case Western Reserve University, United States; {3} New Jersey Institute of Technology, United States; {4} University of Southern California, United States; {5} Intel, United States; {6} AMD, United States; {7} Univ. of Southern California, United States; {8} Colorado State University, United States

16:15 - 17:15

Hardware for Neural Rendering and Geometric Models

Room: Ballroom B

Session Chair(s): Yiyu Shi, University of Notre Dame

Bing Li, University of Siegen

This session tackles real-time 3D neural rendering on resource-constrained platforms through algorithm-hardware co-design. Papers address two key techniques: For 3D Gaussian Splatting, GauPRE introduces a pattern-based rendering engine and LS-Gaussian uses inter-frame continuity for lightweight streaming. For point-based models, SLTarch tames workload imbalance with a dedicated hardware core, while another work details an energy-efficient accelerator for Point Transformer Networks that reduces data access.

16:15

23: SLTarch: Towards Scalable Point-Based Neural Rendering by Taming Workload Imbalance and Memory Irregularity

Xingyang Li {1}, Jie Jiang {1}, Yu Feng {1}, Yiming Gan {2}, Jieru Zhao {1}, Zihan Liu {1}, Jingwen Leng {1}, Minyi Guo {1}

{1} Shanghai Jiao Tong University, China; {2} Institute of Computing Technology, Chinese Academy of Sciences, China

16:30

26: GauPRE: A Pattern-based Rendering Engine for Gaussian Splatting on Edge Device

Yuzheng Lin {1}, Lizhou Wu {1}, Chixiao Chen {2}, Xiaoyang Zeng {3}, Haozhe Zhu {3}

{1} State Key Laboratory of Integrated Chips & Systems, Frontier Institute of Chip and System, Fudan University, China; {2} Fudan University, China; {3} State Key Laboratory of Integrated Chips & Systems, Fudan University, China

Technical Program – Oct. 29

16:45

174: Energy-Efficient Accelerator for Scalable Point Transformer Networks with Reduced Data Access

Hyunsung Yoon {1}, Jehun Lee {2}, Jae-Joon Kim {2}

{1} Pohang University of Science and Technology, Republic of Korea; {2} Seoul National University, Republic of Korea

17:00

1364: No Redundancy, No Stall: Lightweight Streaming 3D Gaussian Splatting for Real-time Rendering

Linye Wei {1}, Jiajun Tang {1}, Fan Fei {1}, Boxin Shi {1}, Runsheng Wang {1}, Meng Li {2}

{1} Peking University, China; {2} Institute for Artificial Intelligence and School of Integrated Circuits, Peking University, China

16:15 - 17:15

Memory-Centric Systems: Design and Optimizations

Room: Barcelona

Session Chair(s): Tathagata Srimani, CMU

John Paul Strachan, RWTH Aachen University

This session surveys circuit-to-system advances that improve reliability, energy efficiency, and on-chip learning capability in emerging memory technologies. Presentations span selector-only memories, a FeFET analog CAMs, a physics-guided generative model for analog CIM non-idealities, and a continuously trainable SOT-MRAM compute-in-memory engine. Together, these works demonstrate how device-aware modeling and adaptive architectures can drastically improve data integrity, search efficiency, and compute throughput in next-generation accelerators.

16:15

1437: Self-Error Detection and Correction Techniques for Reliable and Efficient Selector-Only Memory

Hyunjun Lee, Joon-Sung Yang

Yonsei University, Republic of Korea

16:30

960: FACAM: Design and Optimization of A Compact Energy Efficient FeFET-Based Analog Content Addressable Memory

Jiahao Cai {1}, Ann Franchesca Laguna {2}, Zeyu Yang {1}, Yuxiao Yang {1}, Thomas Kämpfe {3}, Zheyu Yan {1}, Cheng Zhuo {1}, Xunzhao Yin {1}

{1} Zhejiang University, China; {2} De La Salle University, Philippines; {3} Fraunhofer IPMS, Germany

Technical Program – Oct. 29

16:45

484: Towards Accurate Characterization of In-Memory Computing Non-Idealities: A Physics & Data Co-Driven Generative Framework

Jing Kou {1}, Guangyao Wang {1}, Saiya Wang {1}, Yuexi Lv {2}, Liang Zhang {1}, Chenglin Yu {2}, Xinghao Cui {2}, Yulong Liu {1}, Wei Xing {3}, Wang Kang {1}

{1} Beihang University, China; {2} Zhicun Research Lab, China; {3} The University of Sheffield, United Kingdom

17:00

1424: Continuous On-Chip Learning in Neural Networks using SOT-MRAM based CIM Architectures

Anubha Sehgal {1}, Sandeep Soni {2}, Sumit Diware {3}, Alok Shukla {4}, Sourajeet Roy {2}, Rajendra Bishnoi {5}

{1} IIT Roorkee, India; {2} Indian Institute of Technology Roorkee, India; {3} Delft University of Technology, Netherlands; {4} Madan Mohan Malaviya University of Technology, India; {5} Delft University of Technology,, Netherlands

16:15 - 17:15

Multi-Physics Modelling and Advanced Integration

Room: Atlanta

Session Chair(s): Zhuo FENG, Stevens Inst. Tech

Sung Kyu Lim, University of Southern California, US

Delve into the critical intersection of multi-physics modeling and advanced integration technologies that enable next-generation electronic systems. This session addresses the growing complexity of modern packaging and interconnect design through innovative modeling approaches. From physics-informed neural operators predicting warpage in advanced packaging to transformer architectures understanding 3D PCB electromagnetics, these works tackle the multi-domain challenges of wafer-scale integration, signal integrity, and parasitic extraction. Essential for designers working on high-performance computing, AI accelerators, and complex system-in-package solutions.

16:15

590: Advanced Packaging Warpage Modeling with DeepONet-Based Operator Learning

Shao-Yu Lo, Che-Ming Chang, Yao-Wen Chang

National Taiwan University, Taiwan

16:30

1352: MM-Pack: Multi-Mask Co-Design for Ultra-Large Wafer-Scale Package Integration

Shanyi Li, Zhen Zhuang, Siyuan Liang, Bei Yu, Tsung-Yi Ho

The Chinese University of Hong Kong, Hong Kong Special Administrative Region of China

Technical Program – Oct. 29

16:45

358: PCBFormer: Understanding 3D Structure of Real World PCB Traces for S-Parameter Prediction

Taejin Paik {1}, Jaemin Park {2}, Daniel Jung {3}, taehee kim {4}, Doyun Kim {2}
{1} Samsung Electronics, AI Center, Republic of Korea; {2} AI center, Samsung Electronics, Republic of Korea; {3} Samsung Electronics, Republic of Korea; {4} samsung, Republic of Korea

17:00

1142: Capacitance Extraction via Machine Learning with Application to Interconnect Geometry Exploration

Cheng-Yu Tsai {1}, Suwan Kim {2}, Sung Kyu Lim {3}
{1} Georgia Institute of Technology, United States; {2} Samsung Electronics, Republic of Korea; {3} Georgia Tech, United States

16:15 - 17:15

Special Session: Bridging the Gap: Design Automation Meets Microfluidics

Room: Sydney

Session Chair(s): Laurens Spoelstra, University of Twente

Microfluidic devices, or Labs-on-a-Chip (LOCs), have become essential tools for biochemical and medical experimentation. Recent innovations—such as organ-on-chip systems, pandemic-driven diagnostics, and ISO standards—have significantly advanced the field but also introduced new design challenges. These include precise component placement, fluid routing, and process control, which mirror complexities in conventional circuit design. While the design automation community has the expertise to address these issues, its impact on microfluidics remains limited. This special session aims to bridge the gap between design automation and microfluidics by bringing together experts from both fields. It will present practical design challenges, introduce a software framework tailored to microfluidic needs, and showcase a successful Lab-on-a-Chip implementation using adapted automation tools. The session will offer a comprehensive overview from specification to fabrication and foster dialogue to advance design automation for microfluidics.

16:15

20003: Invited Paper: Connecting the D.O.T.S.: Design of fluidic circuit boards for multi-OoC platforms using CAD Tools for Standardization

Laurens Spoelstra, Marlize Kramer, Jasper Rietveld, Joshua Loessberg-Zahl, Loes Segerink
University of Twente, Netherlands

16:30

20004: Invited Paper: The Munich Microfluidic Toolkit: Design Automation and Simulation Tools for Microfluidic Devices

Robert Wille {1}, Philipp Ebner {2}, Maria Emmerich {1}, Michel Takken {1}
{1} Technical University of Munich, Germany; {2} Johannes Kepler University, Austria

Technical Program – Oct. 29

16:45

20005: Invited Paper: Liquid Computing: Towards Programmable Microfluidics

Luca Pezzarossa, Joel August Vest Madsen, Alexander Marc Collignon, Jan Madsen

Technical University of Denmark, Denmark

18:30 – 21:30

ACM SIGDA Dinner

Hofbräuhaus München Festsaal

Technical Program – Oct. 30

7:30 – 8:00

Registration

Room: Upstairs Foyer

8:00 – 16:30

Workshop on Post-Quantum Cryptography Resilience, Verification, and Secure Design Automation (WPQC)

Room: Sydney

The rapid transition toward Post-Quantum Cryptography (PQC) has prompted a global effort to redefine digital security in anticipation of quantum-capable adversaries. While PQC algorithms are mathematically robust, their secure and efficient integration into hardware systems introduces new vulnerabilities and engineering challenges—especially in areas such as side-channel resistance, hardware/software co-design, and formal verification.

The Workshop on Post-Quantum Cryptography and Secure Hardware (WPQC) aims to address these challenges by serving as a dedicated forum for researchers and practitioners in hardware security, EDA/CAD, computer architecture, and cryptography.

The workshop provides a platform to explore the intersection of PQC and hardware/system design, highlighting recent advancements in secure accelerator development, design automation, system integration, and emerging threat models. By bringing together experts from academia, industry, and standardization bodies, WPQC seeks to accelerate collaborative research efforts and foster a shared vision for building resilient, verifiable, and efficient post-quantum systems.

8:00 – 16:30

Top Picks in Hardware and Embedded Security

Room: Barcelona

Top Picks Workshop creates a venue to showcase the best and high impact recently published works in the area of hardware and embedded security. These works will be selected from conference papers that have appeared in leading hardware security conferences including but not limited to DAC, ICCAD, DATE, ASPDAC, HOST, Asian HOST, GLSVLSI, VLSI Design, CHES, ETS, VTS, ITC, S&P, Usenix Security, CCS, NDSS, ISCA, MICRO, ASPLOS, HPCA, HASP, ACSAC, Euro S&P, and Asia CCS. The 8th Top Picks workshop will be collocated with ICCAD 2025. The authors of a short list of papers picked from the submissions are required to present their work at the workshop on October 30, 2025 and are invited to attend ICCAD's networking in-person on October 30, 2025. The presentation including a discussion (Q&A) session, is mandatory.

Technical Program – Oct. 30

8:00 – 16:30

2nd Quantum Computing Applications and Systems (QCAS) Workshop

Room: Athens

Our proposed workshop aims to address emerging challenges and explore innovative solutions in the field of quantum technologies, particularly focusing on quantum computing applications for real-world problems. Quantum technologies are becoming crucial in a variety of scientific domains including chemistry, finance, power systems, etc. Our workshop will bring together researchers, practitioners, and industry experts to exchange ideas, share applications, and discuss the latest advancements in quantum algorithms, error correction, control techniques, and their applications across diverse fields. This event will serve as a platform to showcase cutting-edge research, foster collaboration, and drive innovation in the intersection of CAD, optimization, machine learning, and quantum computing.

8:00 – 16:30

System-Level Interconnect Pathfinding (SLIP)

Room: Rome

SLIP, co-located with ICCAD, brings together researchers and practitioners who have a shared interest in the challenges and futures of system-level interconnect, coming from wide-ranging backgrounds that span system, application, design and technology.

The technical goal of the workshop is to

identify fundamental problems, and,
foster new pathfinding of design, analysis, and optimization of system-level interconnects with emphasis on system-level interconnect modeling and pathfinding, DTCO-enhanced interconnect fabrics, memory and processor communication links, novel dataflow mapping for machine learning, 2.5/3D architectures, and new fabrics for the beyond-Moore era.

Original submissions in the form of regular technical papers, invited sessions (tutorials, panels, special-topic sessions), workshop discussion topics, and posters are welcome. Program content is accepted based on novelty and contributions to the advancement of the field. Authors will keep the copyright of their work and there will be no published proceedings.

Technical Program – Oct. 30

8:00 – 11:30

Foundation Models and EDA

Room: Montreal

The rapid advancement of foundation models has brought powerful new capabilities to Electronic Design Automation (EDA). Unlike traditional task-specific Artificial Intelligence (AI) approaches, foundation models leverage self-supervised learning and large-scale data pre-training to acquire strong generalization abilities. These models can then be efficiently fine-tuned on EDA-specific datasets, enabling a wide range of downstream applications such as hardware code generation, debug and optimization, EDA agents, circuit representation learning and understanding, etc.

This workshop aims to foster collaboration among researchers, engineers, and industry experts to explore the evolving intersection of foundation models and EDA methodologies. By identifying key challenges and exchanging ideas, we seek to inspire comparative analysis between the unique demands of EDA and those of other application domains, and to promote novel solutions that leverage the strengths of foundation models for more accurate, efficient, and scalable design automation. Through talks, paper presentations, and interactive discussions, the workshop will showcase recent progress and explore new ways to apply foundation models in EDA—helping to drive innovation and advance the future of intelligent circuit design.

9:30 – 10:00

Coffee Break

Room: Upstairs Foyer

11:30 – 13:00

Standing Lunch

Room: Foyer Area of Sydney and Atlanta

13:00 – 16:30

Automotive Chiplet Platform: Solutions Enabled by Industry and Academic Collaboration

Room: Montreal

Chiplet-based platform is a pragmatic solution to meet the automotive industry's growing demands for performance, longevity, and cost-efficiency. Such platforms enable modular, scalable integration of diverse processing units (like CPUs, GPUs, and AI accelerators) across different process nodes, improving design flexibility, yield, and reliability, which is critical in safety-driven environments. They also support thermal and power optimization, simplify long-term maintenance and supply chain challenges, and align well with emerging zonal E/E architectures. The purpose of the workshop is to outline the challenges and novel solutions to enable chiplet-based computing platform for automotive domain.

Technical Program – Oct. 30

14:30 – 15:00

Coffee Break

Room: Upstairs Foyer