

Hyunwoo Oh

+1 (949) 571-0953 ✉ hyunwooo@uci.edu 🌐 hyun-woo-oh.github.io 📄 Hyunwoo Oh

EDUCATION

Ph.D. Student in Computer Science

University of California, Irvine (UCI)

- Supervisor: Prof. Mohsen Imani

2024-present

Irvine, CA, USA

M.S. in Computer Science

University of California, Irvine (UCI)

2024-2026

Irvine, CA, USA

M.S. in Electronic Engineering

Seoul National University of Science and Technology (SEOULTECH)

2021-2023

Seoul, Korea

- Supervisor: Prof. Seung Eun Lee

B.S. in Electronic Engineering

Seoul National University of Science and Technology (SEOULTECH)

2012-2021

Seoul, Korea

WORK EXPERIENCE

Research Intern, Microsoft, Redmond, WA, USA

Jun. 2026-Sep. 2026 (expected)

Junior Engineer (Full-time), Hanwha Systems, South Korea

Jan. 2023-Aug. 2024

SELECTED PUBLICATIONS (9 OUT OF 40) [SEE ALL ↓↓]

PolyQ: Codesigning End-to-End Quantization Framework for Scalable Edge CPU LLM Inference.

Hyunwoo Oh, Suyeon Jang, Hanning Chen, KyungIn Nam, Sanggeon Yun, Ryoza Masukawa, and Mohsen Imani.

IEEE/ACM International Conference on Computer-Aided Design (ICCAD), 2026, Accepted.

ExaGEMM: Exploration Framework for CPU-Driven ML Inference via Associative In-Register Computing for Low-Bit GEMM.

Hyunwoo Oh, Suyeon Jang, Hanning Chen, Sanggeon Yun, Ryoza Masukawa, and Mohsen Imani.

IEEE/ACM International Conference on Computer-Aided Design (ICCAD), 2026, Accepted.

TRINE: A Token-Aware, Runtime-Adaptive FPGA Inference Engine for Multimodal AI.

Hyunwoo Oh, Hanning Chen, Sanggeon Yun, Yang Ni, Suyeon Jang, Behnam Khaleghi, Fei Wen and Mohsen Imani.

ACM/IEEE Design Automation Conference (DAC), 2026, Accepted.

TorR: Towards Brain-Inspired Task-Oriented Reasoning via Cache-Oriented Algorithm-Architecture Co-design.

Hyunwoo Oh, SungHeon Jeong, Suyeon Jang, Hanning Chen, Sanggeon Yun, Tamoghno Das and Mohsen Imani.

ACM/IEEE Design Automation Conference (DAC), 2026, Accepted.

T-SAR: A Full-Stack Co-design for CPU-Only Ternary LLM Inference via In-Place SIMD ALU Reorganization.

Hyunwoo Oh, KyungIn Nam, Rajat Bhattachariya, Hanning Chen, Tamoghno Das, Sanggeon Yun, Suyeon Jang, Andrew Ding, Nikil Dutt, and Mohsen Imani.

Design, Automation and Test in Europe Conference (DATE), 2026.

QUILL: An Algorithm-Architecture Co-Design for Cache-Local Deformable Attention.

Hyunwoo Oh, Hanning Chen, Sanggeon Yun, Yang Ni, Wenjun Huang, Tamoghno Das, Suyeon Jang, and Mohsen Imani.

Design, Automation and Test in Europe Conference (DATE), 2026.

DL-Sort: A Hybrid Approach to Scalable Hardware-Accelerated Fully-Streaming Sorting.

Hyun Woo Oh, Joungmin Park, and Seung Eun Lee.

IEEE Transactions on Circuits and Systems II: Express Briefs (TCAS-II), vol. 71, no.5, 2024.

iTaskSense: Task-Oriented Object Detection in Resource-Constrained Environments.

SungHeon Jeong, Hamza Errahmouni Barkam, Hyunwoo Oh, Hanning Chen, Tamoghno Das, Zhen Ye, and Mohsen Imani.

ACM/IEEE Design Automation Conference (DAC), 2025.

LVL CSP: Accelerating Large Vision Language Models via Clustering, Scattering, and Pruning for Reasoning Segmentation.

Hanning Chen, Yang Ni, Wenjun Huang, Hyunwoo Oh, Yezi Liu, Tamoghno Das, and Mohsen Imani.

ACM International Conference on Multimedia (MM), 2025.

PROFESSIONAL SERVICES

Technical Program Committee Member, 40th ACM/SIGAPP Symposium On Applied Computing (SAC 2025), 4 papers reviewed

2024-2025

Reviewer, IEEE Transactions on Circuits and Systems I: Regular Papers, 1 paper reviewed

2024-

Reviewer, IEEE Access, 5 papers reviewed

2023-

RESEARCH EXPERIENCES

Research Intern, Microsoft, Redmond, WA, USA

Jun. 2026 – Present (expected Sep. 2026)

- Conducting research on FPGA-based compute and memory modeling for advanced accelerator architectures.
- Developing RTL-based hardware emulation logic to model compute and memory subsystem behavior in FPGA-based research frameworks.
- Supporting functional and timing-fidelity validation against architecture specifications and system-level performance goals.
- Using tools and methods including Verilog/SystemVerilog, FPGA development workflows, simulation and verification tools, Python, TCL, and computer architecture modeling infrastructure.

Bio-Inspired Architecture and Systems Lab., UC Irvine (PI: Prof. Mohsen Imani)

Sep. 2024 - current

Quantization-Aware SW/HW Co-design for CPU-Only LLM Inference

Mar. 2025 – Present

- Co-designed ternary LLM GEMM/GEMV algorithms, a ternary ISA extension, AVX2 kernels, and SIMD-unit RTL for LLM inference (e.g., BitNet-b1.58).
- Repurposed SIMD register files as ternary LUTs to remove LUT DRAM traffic, achieving **5.6–24.5× lower GEMM latency**, **1.1–86.2× higher GEMV throughput**, and **2.5–4.9× better energy** than Jetson AGX Orin with only **3.2% power / 1.4% area overhead** for the SIMD unit redesign.
- Developed an intra-layer mixed-precision quantization and compilation framework with activation-aware per-channel 2/3/4/8/16-bit weights and AVX2-friendly kernels.
- Achieved effective **3–6 b** scaling on Falcon-H1-3B, Llama2-13B, and Qwen3-32B, reducing perplexity by **2.4–7.8% vs AWQ** and **10.9–17.0% vs GPTQ** at similar latency/energy.
- **Publications: DATE 2026 (C23).**

ASIC Accelerator for Emerging Object Detection Models

Sep. 2024 – Present

- Designed QUILL, a deformable-attention accelerator for DETR-style detection, co-optimizing attention schedules with on-chip memory.
- Achieved up to **7.3× throughput** and **47× energy efficiency vs RTX 4090** at comparable accuracy by reordering queries and prefetching regions to make sparse deformable-attention accesses cache-local in a 28nm fused core.
- Designed a **128×128 8-bit systolic array** with flexible weight/input/output-stationary dataflows and on-chip scheduling, removing SRAM and host tensor-reordering overheads.
- Achieved **30.3 TOPS @ 1.85 GHz (28nm)** with up to **3.5× speedup** and **40% lower energy** vs GPU baselines.
- **Publications: DAC 2026 (C25), DATE 2026 (C22), DAC 2025 (C15).**

FPGA Accelerator for Multimodal AI Workloads

Dec. 2024 – Present

- Designed an FPGA accelerator and compiler for multimodal AI (ViT, GNN, CNN, NLP) under tight latency and resource budgets.
- Implemented mode-switchable compute engines, scalable top-k token pruning, and dependency-aware offloading to support diverse workloads without bitstream reconfiguration.
- Developed a Chisel-based RTL generator targeting Xilinx Alveo U50 and ZCU104, achieving up to **22.6× lower latency vs RTX 4090** and **6.9× vs Jetson Orin Nano**, outperforming prior FPGA accelerators.
- **Publications: DAC 2026 (C25), DATE 2026 (C21), FCCM 2025 (C12), ISLPED 2025 (C14).**

Hardware Performance Analyses for Hyperdimensional Computing Optimization Methods

Apr. 2025 – Present

- Analyzed hardware performance of compression algorithms on HDC classifiers via logarithmic class-axis reduction, quantifying robustness and efficiency gains under hardware constraints.
- Decomposition representation for Extreme-memory HDC: Measured accuracy/efficiency trade-offs for constrained execution.
- Performance analysis on scalable symbolic reasoning using matrix-based brain-inspired representations with vector-space acceleration.
- Evaluated the compute/memory benefits of robust, efficient direction-of-arrival (DoA) estimation using HDC.
- **Publications: DATE 2026 (C20), DATE 2026 (C19), DATE 2026 (C18), ICASSP 2026 (C16).**

Algorithm and Model Studies for Hardware-Efficient Machine Learning

Jun. 2025 – Present

- Contributed to the engineering and evaluation of a training-free token pruning framework for improving the computational efficiency of large vision-language models in reasoning-based segmentation.
- Supported benchmarking across datasets and models, helping validate inference-cost reductions while preserving segmentation performance.
- **Publication: MM 2025 (C14).**

Automated Hardware Generation and LLM-Aided Design

Apr. 2025 – Present

- Developed an agentic, memory-aware framework that uses LLM assistance to generate hardware more reliably under resource constraints.
- **Publications: VLSID 2026 (C17).**

Junior Engineer (Full-time), Core H/W Team, Hanwha Systems, South Korea

Jan. 2023 - Aug. 2024

High-Resolution Thermal Image Processor for 1280×1024 IR Cameras

Jan. 2023 – Aug. 2024

- Built an end-to-end thermal image pipeline (NUC + CLAHE) on Zynq SoC FPGAs with PL accelerators and ARM cores sharing DDR via AXI.
- Achieved real-time 640×480 @ 60 FPS on XC7Z020 using <40% LUT/BRAM/DSP; extended to Zynq Ultrascale+ MPSoC for 1280×1024 @ 60 FPS in a fully on-device sensor-to-display pipeline.

- **Publications:** DSD 2023 (C7), DSD 2023 (C6).

Real-Time Thermal Imaging on Heterogeneous SoC

Jan. 2023 – Aug. 2024

- Developed NUC, contrast enhancement, and noise filtering on Cortex-M4, dual TI C66x DSPs, and vision engine under memory/power limits.
- Hand-optimized fixed-point SIMD/VLIW kernels and cache-aware task scheduling to reach 640×480 @ 60 FPS with <2.2 W system power and 57.5% peak core load.
- **Publications:** RTCSA 2024 (C11).

Computer Architecture Lab., SEOULTECH, South Korea (PI: Prof. Seung Eun Lee)

Dec. 2019 – Feb. 2023

Sorting Accelerator for Resource-Constrained Platforms

Mar. 2022 – Dec. 2022

- Designed a scalable sorting accelerator that balances fully-streaming capability and hardware resource usage.
- **Publications:** TCAS-II (J9).

Processor Architecture and Number Format Studies for Edge Machine Learning

Jun. 2020 – Dec. 2022

- Designed a lightweight RISC core supporting efficient IEEE-754 to Posit migration.
- Designed a lightweight RISC core with scalable *k*-NN acceleration support.
- Designed an RTL generator for a reconfigurable embedded AI accelerator supporting *k*-NN and RBF-NN.
- Researched SW applications of the *k*-NN accelerator module.
- **Publications:** ISLPED 2023 (C8), IEEE Access (J7), ISOCC 2022 (C5), Micromachines (J3), Micromachines (J2), ICCE 2021 (C2), JICCE (J5), ICFICE 2022 (C3), ISOCC2022 (C5), ISOCC 2020 (C1).

Domain-Specific Accelerators for Respiratory Medical Device

Jan. 2020 – Mar. 2021

- Designed a 2D graphics accelerator optimized for graph visualization tasks in respiratory medical devices.
- Developed a SW stack for Lempel-Ziv 77 (LZ77) lossless decompression accelerator: ① Modeling LZ77 algorithm to design hardware. ② Developing a SW stack for data pre-processing and evaluation.
- **Publications:** Electronics (J4), Micromachines (J1).

Other projects: Configurable JTAG TAP RTL generator, LIN Controller IP [ICCE 2022 (C4)].

Participated in several ASIC design projects using Synopsys EDA tools. [\[See list ↗\]](#)

TEACHING EXPERIENCE

Teaching Assistant for “Computer Architecture”, SEOULTECH

Fall 2021

Teaching Assistant for “Digital System Design”, SEOULTECH

Spring 2021

AWARDS AND HONORS

Computer Science Department Research Fellowship	University of California, Irvine	\$ 18.7K	2024
Discretionary Merit Award	University of California, Irvine	\$ 2.5K	2024
Future Talent Scholarship	Seoul National University of Science and Technology	\$ 5.5K	2021-2022
Academic Scholarship	Seoul National University of Science and Technology	\$ 0.4K	2021
President of the Institute of Semiconductor Engineers Award	21st Korea Semiconductor Design Contest	\$ 0.9K	2020

TECHNICAL SKILLS

RTL Design	HDL & HLS Simulation	SystemVerilog, Kanagawa, Chisel Verilator, ModelSim
System Modeling & Simulation	CPU	gem5
Cell-Based ASIC Design	Synopsys EDA Tools	Design Compiler (Synthesis), IC Compiler I/II (Layout) VCS (Simulation), Verdi (Analysis), Formality (Validation) StarRCXT (Parasitic Extraction), PrimeTime (STA)
FPGA Design	Cadence EDA Tools Xilinx FPGA Tools Intel FPGA Tools	Virtuoso Layout Suite (Layout), Calibre DRC/LVS (Physical/Layout Verification) Vivado, Vitis Quartus II/Prime, Nios II EDS
Computer Programming	Languages ML Frameworks OS Development	C/C++, Python PyTorch, CUDA, Numpy FreeRTOS, TI Vision SDK, PetaLinux

TRAINING

Design of High-speed Memory Interface, IDEC	2022.12.09
Cell-based Chip Design Flow for Samsung 28nm Process, IDEC	2021.11.01-11.05
[Synopsys] Block-level Auto P&R utilizing IC Compiler II, IDEC	2021.10.19-10.21
Cell-based Chip Design Flow, IDEC	2021.07.05-07.09
[Infineon] Automotive Semiconductor Expert Training - Basic Course, KSIA	2021.06.30-07.02
Cell-based Chip Design Flow, IDEC	2020.08.10-08.14

ALL PUBLICATIONS [\[Go up ↑\]](#)

Conferences

- C30. **PolyQ: Codesigning End-to-End Quantization Framework for Scalable Edge CPU LLM Inference.**
Hyunwoo Oh, Suyeon Jang, Hanning Chen, KyungIn Nam, Sanggeon Yun, Ryoza Masukawa, and Mohsen Imani.
IEEE/ACM International Conference on Computer-Aided Design (ICCAD), 2026, Accepted.
- C29. **ExaGEMM: Exploration Framework for CPU-Driven ML Inference via Associative In-Register Computing for Low-Bit GEMM.**
Hyunwoo Oh, Suyeon Jang, Hanning Chen, Sanggeon Yun, Ryoza Masukawa, and Mohsen Imani.
IEEE/ACM International Conference on Computer-Aided Design (ICCAD), 2026, Accepted.
- C28. **Qubit-Efficient Quantum Search for Hyperdimensional Decomposition via Logarithmic Encoding.**
Sanggeon Yun, **Hyunwoo Oh**, Ryoza Masukawa, Raheeb Hassan, and Mohsen Imani.
IEEE/ACM International Conference on Computer-Aided Design (ICCAD), 2026, Accepted.
- C27. **TRINE: A Token-Aware, Runtime-Adaptive FPGA Inference Engine for Multimodal AI.**
Hyunwoo Oh, Hanning Chen, Sanggeon Yun, Yang Ni, Suyeon Jang, Behnam Khaleghi, Fei Wen, and Mohsen Imani.
ACM/IEEE Design Automation Conference (DAC), 2026, Accepted.
- C26. **TorR: Towards Brain-Inspired Task-Oriented Reasoning via Cache-Oriented Algorithm-Architecture Co-design.**
Hyunwoo Oh, SungHeon Jeong, Suyeon Jang, Hanning Chen, Sanggeon Yun, Tamoghno Das, and Mohsen Imani.
ACM/IEEE Design Automation Conference (DAC), 2026, Accepted.
- C25. **T-SAR: A Full-Stack Co-design for CPU-Only Ternary LLM Inference via In-Place SIMD ALU Reorganization.**
Hyunwoo Oh, KyungIn Nam, Rajat Bhattacharjya, Hanning Chen, Tamoghno Das, Sanggeon Yun, Suyeon Jang, Andrew Ding, Nikil Dutt, and Mohsen Imani.
Design, Automation and Test in Europe Conference (DATE), 2026.
- C24. **QUILL: An Algorithm-Architecture Co-Design for Cache-Local Deformable Attention.**
Hyunwoo Oh, Hanning Chen, Sanggeon Yun, Yang Ni, Wenjun Huang, Tamoghno Das, Suyeon Jang, and Mohsen Imani.
Design, Automation and Test in Europe Conference (DATE), 2026.
- C23. **RIFT: A Single-Bitstream, Runtime-Adaptive FPGA-Based Accelerator for Multimodal AI.**
Hyunwoo Oh, Hanning Chen, Sanggeon Yun, Yang Ni, Suyeon Jang, Behnam Khaleghi, Fei Wen, and Mohsen Imani.
Design, Automation and Test in Europe Conference (DATE), 2026.
- C22. **LogHD: Robust Compression of Hyperdimensional Classifiers via Logarithmic Class-Axis Reduction.**
Sanggeon Yun, **Hyunwoo Oh**, Ryoza Masukawa, Pietro Mercati, Nathaniel D Bastian, and Mohsen Imani.
Design, Automation and Test in Europe Conference (DATE), 2026.
- C21. **DecoHD: Decomposed Hyperdimensional Classification under Extreme Memory Budgets.**
Sanggeon Yun, **Hyunwoo Oh**, Ryoza Masukawa, and Mohsen Imani.
Design, Automation and Test in Europe Conference (DATE), 2026.
- C20. **Scalable Symbolic Reasoning with Matrix-Based Brain-Inspired Representations and Vector-Space Acceleration.**
William Youngwoo Chung, **Hyunwoo Oh**, Hamza Errahmouni Barkam, Calvin Yeung, and Mohsen Imani.
Design, Automation and Test in Europe Conference (DATE), 2026.
- C19. **Vector-Space Projection and Unified Complex-Valued Acceleration for Scaling Matrix-Based Brain-Inspired Representations.**
William Youngwoo Chung, **Hyunwoo Oh**, Calvin Yeung, Hansen Jin Lillemark, Hamza Errahmouni Barkam and Mohsen Imani.
ACM/IEEE International Symposium on Low Power Electronics and Design (ISLPED), 2026.
- C18. **FusionSense: Tri-Stage Near-Sensor Learning for Runtime-Adaptive Multimodal Edge Intelligence.**
Sanggeon Yun, Ryoza Masukawa, Minhyoung Na, **Hyunwoo Oh**, Yoshiki Yamaguchi, Wenjun Huang, SungHeon Jeong and Mohsen Imani.
ACM/IEEE International Symposium on Low Power Electronics and Design (ISLPED), 2026.
- C17. **AAMLA: An Autonomous Agentic Framework for Memory-Aware LLM-Aided Hardware Generation.**
Rajat Bhattacharjya, Juhee Sung, Hangyeol Jung, **Hyunwoo Oh**, Arnab Sarkar, Mohsen Imani and Nikil Dutt.
International Conference On VLSI Design (VLSID), 2026.
- C16. **HYPERDOA: Robust and Efficient DoA Estimation using Hyperdimensional Computing.**
Rajat Bhattacharjya, Woohyeok Park, Arnab Sarkar, **Hyunwoo Oh**, Mohsen Imani, and Nikil Dutt.
IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), 2026.

C15. iTaskSense: Task-Oriented Object Detection in Resource-Constrained Environments.

SungHeon Jeong, Hamza Errahmouni Barkam, **Hyunwoo Oh**, Hanning Chen, Tamoghno Das, Zhen Ye, and Mohsen Imani.
ACM/IEEE Design Automation Conference (DAC), 2025.

C14. LVLm CSP: Accelerating Large Vision Language Models via Clustering, Scattering, and Pruning for Reasoning Segmentation.

Hanning Chen, Yang Ni, Wenjun Huang, **Hyunwoo Oh**, Yezi Liu, Tamoghno Das, and Mohsen Imani.
ACM International Conference on Multimedia (MM), 2025.

C13. Revisiting Reconfigurable Acceleration of Vision Transformer with Patch Pruning.

Hanning Chen, Yang Ni, Wenjun Huang, **Hyun Woo Oh**, Tamoghno Das, Fei Wen, and Mohsen Imani.
ACM/IEEE International Symposium on Low Power Electronics and Design (ISLPED), 2025.

C12. A Multimodal AI Acceleration with Dynamic Pruning and Run-time Configuration.

Hyun Woo Oh, Hanning Chen, Sanggeon Yun, Yang Ni, Behnam Khaleghi, Fei Wen, and Mohsen Imani.
IEEE International Symposium on Field-Programmable Custom Computing Machines (FCCM), 2025, extended abstract.

C11. A Compact Real-Time Thermal Imaging System Based on Heterogeneous System-on-Chip.

Hyun Woo Oh, Cheol-Ho Choi, Jeong Woo Cha, Hyunmin Choi, Jung-Ho Shin, and Joon Hwan Han.
IEEE International Conference on Embedded and Real-Time Computing Systems and Applications (RTCSA), 2024.

C10. Fast Object Detection Algorithm using Edge-based Operation Skip Scheme with Viola-Jones Method.

Cheol-Ho Choi, Joonhwan Han, Jeongwoo Cha, Jungho Shin, and **Hyun Woo Oh**.
IEEE International Conference on Artificial Intelligence Circuits and Systems (AICAS), 2024.

C9. Algorithm for LWIR Thermal Imaging Camera with Minimal Mechanical Shutter Utilization.

Taehyun Kim, Joonhwan Han, Jeongwoo Cha, Hyunmin Choi, Jungho Shin, Eunchong Kim, **Hyun Woo Oh**, Cheol-Ho Choi, Seongtaek Hong, and Taehyung Kim.
IEEE International Conference on Consumer Electronics-Asia (ICCE-Asia), 2024.

C8. RF2P: A Lightweight RISC Processor Optimized for Rapid Migration from IEEE-754 to Posit.

Hyun Woo Oh, Seongmo An, Won Sik Jeong, and Seung Eun Lee.
ACM/IEEE International Symposium on Low Power Electronics and Design (ISLPED), 2023.

C7. An SoC FPGA-based Integrated Real-time Image Processor for Uncooled Infrared Focal Plane Array.

Hyun Woo Oh, Cheol-Ho Choi, Jeong Woo Cha, Hyunmin Choi, Joon Hwan Han, Jung-Ho Shin.
Euromicro Conference on Digital System Design (DSD), 2023.

C6. Disparity Refinement Processor Architecture utilizing Horizontal and Vertical Characteristics for Stereo Vision Systems.

Cheol-Ho Choi, and **Hyun Woo Oh**.
Euromicro Conference on Digital System Design (DSD), 2023.

C5. Evaluation of Posit Arithmetic on Machine Learning based on Approximate Exponential Functions.

Hyun Woo Oh, Won Sik Jeong, and Seung Eun Lee.
International SoC Design Conference (ISOCC), 2022.

C4. A Local Interconnect Network Controller for Resource-Constrained Automotive Devices.

Kwonneung Cho, **Hyun Woo Oh**, Jeongeun Kim, Young Woo Jeong, and Seung Eun Lee.
IEEE International Conference on Consumer Electronics (ICCE), 2022.

C3. Intelligent Transportation System based on an Edge AI.

Young Woo Jeong, **Hyun Woo Oh**, Su Yeon Jang, and Seung Eun Lee.
International Conference on Future Information & Communication Engineering (ICFICE), 2022.

C2. Vision-based Parking Occupation Detecting with Embedded AI Processor.

Kwonneung Cho, **Hyun Woo Oh**, and Seung Eun Lee.
IEEE International Conference on Consumer Electronics (ICCE), 2021.

C1. Design of 32-bit Processor for Embedded Systems.

Hyun Woo Oh, Kwon Neung Cho, and Seung Eun Lee.
International SoC Design Conference (ISOCC), 2021.

Journals

J10. EOS: Edge-Based Operation Skip Scheme for Real-Time Object Detection Using Viola-Jones Classifier.

Cheol-Ho Choi, Joonhwan Han, **Hyun Woo Oh**, Jeongwoo Cha, and Jungho Shin.
Electronics, Vol. 14, No. 2, 2025.

J9. DL-Sort: A Hybrid Approach to Scalable Hardware-Accelerated Fully-Streaming Sorting.

Hyun Woo Oh, Joungmin Park, and Seung Eun Lee.
IEEE Transactions on Circuits and Systems II: Express Briefs, vol. 71, no.5, 2024.

J8. Contrast Enhancement Method using Region-based Dynamic Clipping Technique for LWIR-based Thermal Camera of Night Vision Systems.

Cheol-Ho Choi, Joonhwan Han, Jeongwoo Cha, Hyunmin Choi, Jungho Shin, Taehyun Kim, and **Hyun Woo Oh**.
Sensors, Vol. 24, No. 12, 2024.

J7. The Design of Optimized RISC Processor for Edge Artificial Intelligence Based on Custom Instruction Set Extension.

Hyun Woo Oh, and Seung Eun Lee.

IEEE Access, Vol. 11, 2023.

J6. Cell-Based Refinement Processor Utilizing Disparity Characteristics of Road Environment for SGM-based Stereo Vision Systems.

Cheol-Ho Choi, Hyun Woo Oh, JoonHwan Han, and Jungho Shin.

IEEE Access, vol. 11, 2023.

J5. An Edge AI Device based Intelligent Transportation System.

Youngwoo Jeong, Hyun Woo Oh, Soohee Kim, and Seung Eun Lee.

Journal of Information and Communication Convergence Engineering, Vol. 20, No. 3, 2022.

J4. The Design of a 2D Graphics Accelerator for Embedded Systems.

Hyun Woo Oh, Ji Kwang Kim, Gwan Beom Hwang, and Seung Eun Lee.

Electronics, Vol. 10, No. 4, 2021.

J3. A Multi-Core Controller for an Embedded AI System Supporting Parallel Recognition.

Suyeon Jang, Hyun Woo Oh, Young Hyun Yoon, Dong Hyun Hwang, Won Sik Jeong, and Seung Eun Lee.

Micromachines, Vol. 12, No. 8, 2021.

J2. ASimOV: A Framework for Simulation and Optimization of an Embedded AI Accelerator.

Dong Hyun Hwang, Chang Yeop Han, Hyun Woo Oh, and Seung Eun Lee.

Micromachines, Vol. 12, No. 7, 2021.

J1. Lossless Decompression Accelerator for Embedded Processor with GUI.

Gwan Beom Hwang, Kwon Neung Cho, Chang Yeop Han, Hyun Woo Oh, Young Hyun Yoon, and Seung Eun Lee.

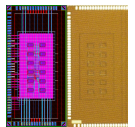
Micromachines, Vol. 12, No. 2, 2021.

CHIP DESIGNS [GO UP ↑↑]

A RISC-V Processor Supporting AMBA AXI Protocol for Embedded Systems

Jul. 2022

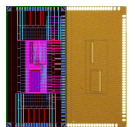
- Designer: Won Sik Jeong, Sun Beom Kwon, Hyun Woo Oh, Jeongeun Kim
- Technology: Samsung 28nm RFCMOS (1-poly 8-metal)
- Role: RTL Verification



Robot-Specific Processor for Autonomous Driving

Jul. 2022

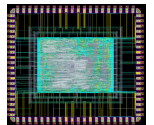
- Designer: Youngwoo Jeong, Yue Ri Jeong, Hyun Woo Oh, Kwang Hyun Go
- Technology: Samsung 28nm RFCMOS (1-poly 8-metal)
- Role: System Verification



In-Vehicle Network Processor based on Cortex-M0

Mar. 2022

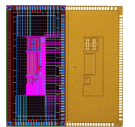
- Designer: Kwon Neung Cho, Jeong Eun Kim, Hyun Woo Oh
- Technology: TSMC 180nm RFCMOS (1-poly 6-metal)
- Role: System Verification SW Dev., RTL Verification, Pre/Post-Layout Simulation



A Programmable Embedded AI Processor with Cortex-M0

Jul. 2021

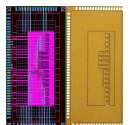
- Designer: Kwon Neung Cho, Young Woo Jeong, Hyun Woo Oh, Chang Yeop Han
- Technology: Samsung 28nm RFCMOS (1-poly 8-metal)
- Role: RTL Subblock Design



32-bit Processor with Posit Arithmetic Coprocessor for Embedded Systems

Jul. 2021

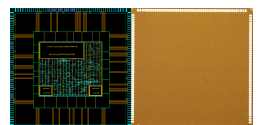
- Designer: Hyun Woo Oh, Jeong Eun Kim, Do Young Choi, Kwang Hyun Go
- Technology: Samsung 28nm RFCMOS (1-poly 8-metal)
- Role: RTL Design & Verification, ASIC Design Front-end/Back-end, Firmware, PCB Design & Chip Test



Implementation of Lossless Decompression Accelerator Based on Inflate Algorithm

Sep. 2020

- Designer: Gwan Beom Hwang, Do Young Choi, Hyun Woo Oh, Chang Yeop Han
- Technology: Samsung 65nm RFCMOS (1-poly 8-metal)
- Role: System Verification SW Dev., PCB Design & Chip Test



Communication System with Simple and Fast Communication Error Check Code Based on CRC

Jun. 2020

- Designer: Chang Yeo Hanp, Kwon Neung Cho, **Hyun Woo Oh**
- Technology: Magnachip Hynix 0.18um CMOS
- Role: RTL Subblock Design

