

2nd Workshop on Memory-Centric Computing: Processing-Near-Memory - Part II

Dr. Geraldo F. Oliveira

<https://geraldofojunior.github.io>

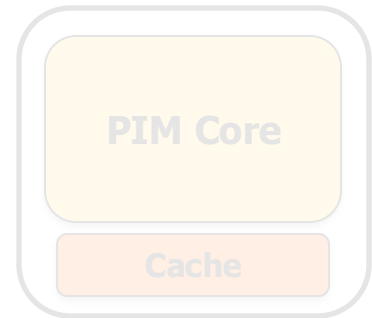
ICS 2025

08 June 2025

Possible PNM Designs

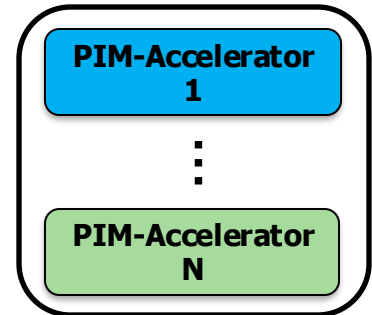
General-purpose programmable cores

- ❑ Wimpy cores (possibility of running any workload)
- ❑ E.g. from academia: Tesseract PIM for Graph Processing
- ❑ E.g. from industry: UPMEM PIM



Fixed-function units

- ❑ Hardware/software co-designed PIM for efficiency
- ❑ E.g. from academia: Mensa for NN Edge Inference
- ❑ E.g. from industry: Samsung HBM-PIM, SK hynix AiM



Reconfigurable architectures

- ❑ PNM cores coupled with FPGAs, CGRA
- ❑ E.g. from academia: NERO for Weather Prediction
- ❑ E.g. from industry: Samsung AxDIMM



Samsung Function-in-Memory DRAM (2021)



Samsung Develops Industry's First High Bandwidth Memory with AI Processing Power

Korea on February 17, 2021

Audio



Share



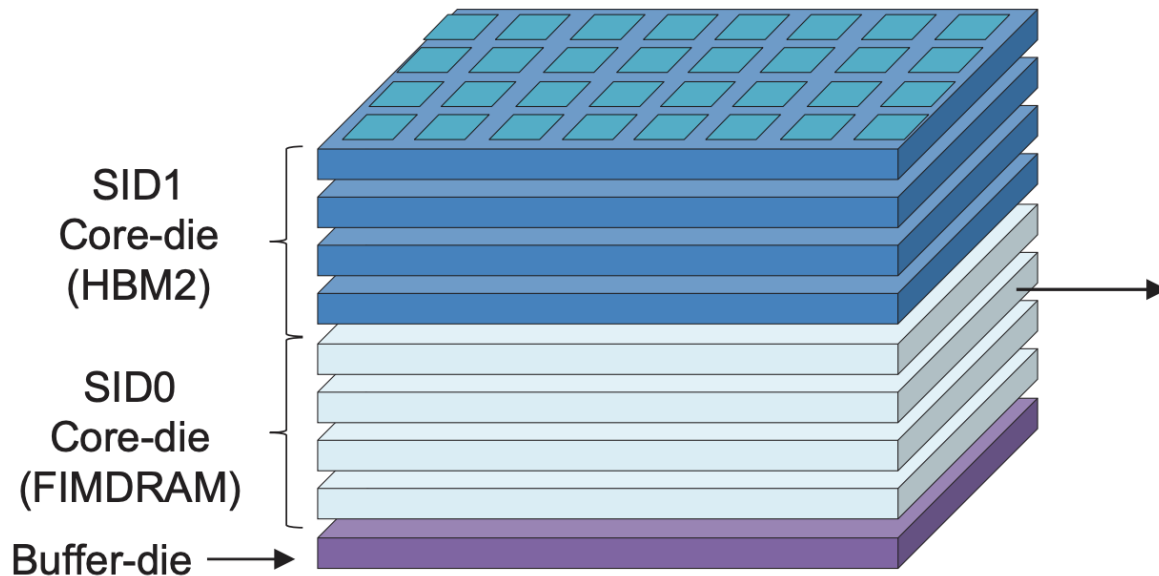
The new architecture will deliver over twice the system performance and reduce energy consumption by more than 70%

Samsung Electronics, the world leader in advanced memory technology, today announced that it has developed the industry's first High Bandwidth Memory (HBM) integrated with artificial intelligence (AI) processing power – the HBM-PIM. **The new processing-in-memory (PIM) architecture brings powerful AI computing capabilities inside high-performance memory, to accelerate large-scale processing in data centers, high performance computing (HPC) systems and AI-enabled mobile applications.**

Kwangil Park, senior vice president of Memory Product Planning at Samsung Electronics stated, "Our groundbreaking HBM-PIM is the industry's first programmable PIM solution tailored for diverse AI-driven workloads such as HPC, training and inference. We plan to build upon this breakthrough by further collaborating with AI solution providers for even more advanced PIM-powered applications."

Samsung Function-in-Memory DRAM (2021)

■ FIMDRAM based on HBM2



[3D Chip Structure of HBM with FIMDRAM]

Chip Specification

128DQ / 8CH / 16 banks / BL4

32 PCU blocks (1 FIM block/2 banks)

1.2 TFLOPS (4H)

**FP16 ADD /
Multiply (MUL) /
Multiply-Accumulate (MAC) /
Multiply-and- Add (MAD)**

ISSCC 2021 / SESSION 25 / DRAM / 25.4

25.4 A 20nm 6GB Function-In-Memory DRAM, Based on HBM2 with a 1.2TFLOPS Programmable Computing Unit Using Bank-Level Parallelism, for Machine Learning Applications

Young-Cheon Kwon¹, Suk Han Lee¹, Jaehoon Lee¹, Sang-Hyuk Kwon¹, Je Min Ryu¹, Jong-Pil Son¹, Seongil O¹, Hak-Soo Yu¹, Haesuk Lee¹, Soo Young Kim¹, Youngmin Cho¹, Jin Guk Kim¹, Jongyoon Choi¹, Hyun-Sung Shin¹, Jin Kim¹, BengSeng Phuah¹, HyoungMin Kim¹, Myeong Jun Song¹, Ahn Choi¹, Daeho Kim¹, SooYoung Kim¹, Eun-Bong Kim¹, David Wang², Shinhaeng Kang¹, Yuhwan Ro³, Seungwoo Seo³, JoonHo Song³, Jaeyoun Youn¹, Kyomin Sohn¹, Nam Sung Kim¹

¹Samsung Electronics, Hwaseong, Korea

²Samsung Electronics, San Jose, CA

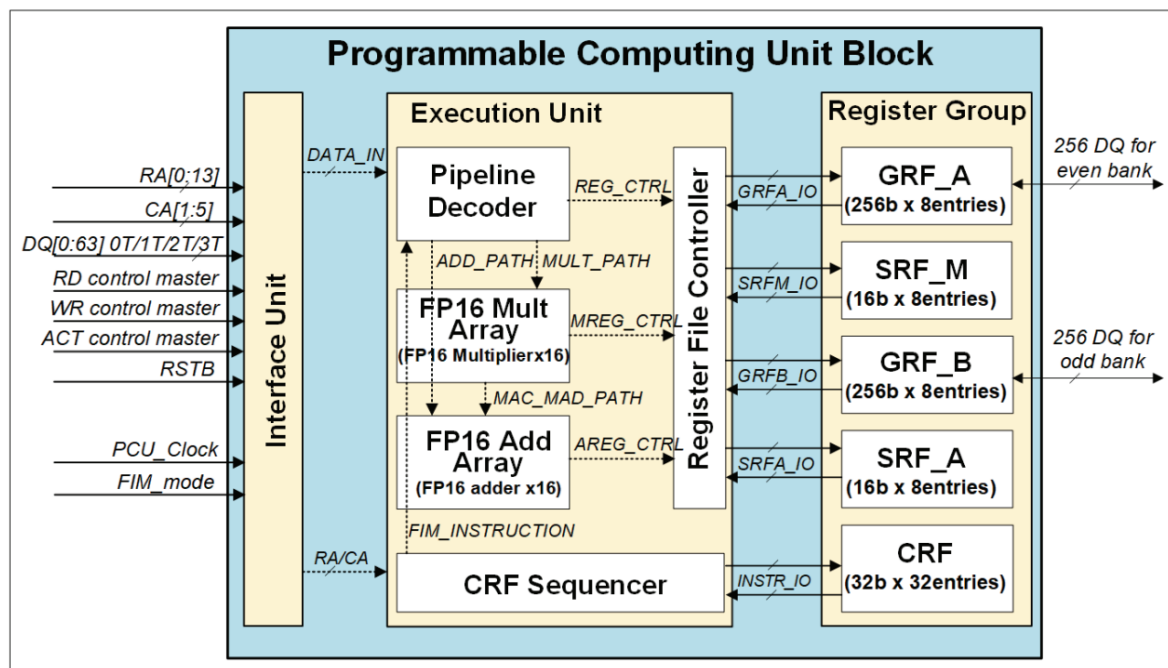
³Samsung Electronics, Suwon, Korea

Samsung Function-in-Memory DRAM (2021)

Programmable Computing Unit

■ Configuration of PCU block

- Interface unit to control data flow
- Execution unit to perform operations
- Register group
 - 32 entries of CRF for instruction memory
 - 16 GRF for weight and accumulation
 - 16 SRF to store constants for MAC operations



[Block diagram of PCU in FIMDRAM]

ISSCC 2021 / SESSION 25 / DRAM / 25.4

25.4 A 20nm 6Gb Function-in-Memory DRAM, Based on HBM2 with a 1.2TFLOPS Programmable Computing Unit Using Bank-Level Parallelism, for Machine Learning Applications

Young-Cheon Kwon¹, Suk Han Lee¹, Jaehoon Lee¹, Sang-Hyuk Kwon¹, Je Min Ryu¹, Jong-Pil Son¹, Seongil O¹, Hak-Soo Yu¹, Haesuk Lee¹, Soo Young Kim¹, Youngmin Cho¹, Jin Guk Kim¹, Jongyoon Choi¹, Hyun-Sung Shin¹, Jin Kim¹, BengSeng Phuah¹, HyungMin Kim¹, Myeong Jun Song¹, Ahn Choi¹, Daeho Kim¹, Sooyoung Kim¹, Eun-Bong Kim¹, David Wang², Shinhaeng Kang¹, Yuhwan Ro², Seungwoo Seo², JoonHo Song¹, Jaeyoun Youn¹, Kyomin Sohn¹, Nam Sung Kim¹

¹Samsung Electronics, Hwaseong, Korea
²Samsung Electronics, San Jose, CA
³Samsung Electronics, Suwon, Korea

Samsung Function-in-Memory DRAM (2021)

[Available instruction list for FIM operation]

Type	CMD	Description
Floating Point	ADD	FP16 addition
	MUL	FP16 multiplication
	MAC	FP16 multiply-accumulate
	MAD	FP16 multiply and add
Data Path	MOVE	Load or store data
	FILL	Copy data from bank to GRFs
Control Path	NOP	Do nothing
	JUMP	Jump instruction
	EXIT	Exit instruction

ISSCC 2021 / SESSION 25 / DRAM / 25.4

25.4 A 20nm 6GB Function-in-Memory DRAM, Based on HBM2 with a 1.2TFLOPS Programmable Computing Unit Using Bank-Level Parallelism, for Machine Learning Applications

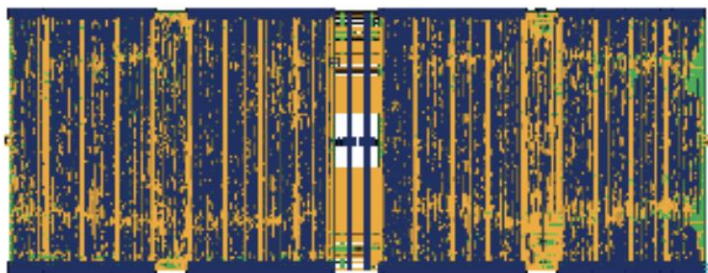
Young-Cheon Kwon¹, Suk Han Lee¹, Jaehoon Lee¹, Sang-Hyuk Kwon¹, Je Min Ryu¹, Jong-Pil Son¹, Seongil O¹, Hak-Soo Yu¹, Haesuk Lee¹, Soo Young Kim¹, Youngmin Cho¹, Jin Guk Kim¹, Jongyoon Choi¹, Hyun-Sung Shin¹, Jin Kim¹, BengSeng Phuah¹, HyounMin Kim¹, Myeong Jun Song¹, Ahn Choi¹, Daeho Kim¹, SooYoung Kim¹, Eun-Bong Kim¹, David Wang², Shinhaeng Kang¹, Yuhwan Ro², Seungwoo Seo², JoonHo Song¹, Jaeyoun Youn¹, Kyomin Sohn¹, Nam Sung Kim¹

¹Samsung Electronics, Hwaseong, Korea
²Samsung Electronics, San Jose, CA
³Samsung Electronics, Suwon, Korea

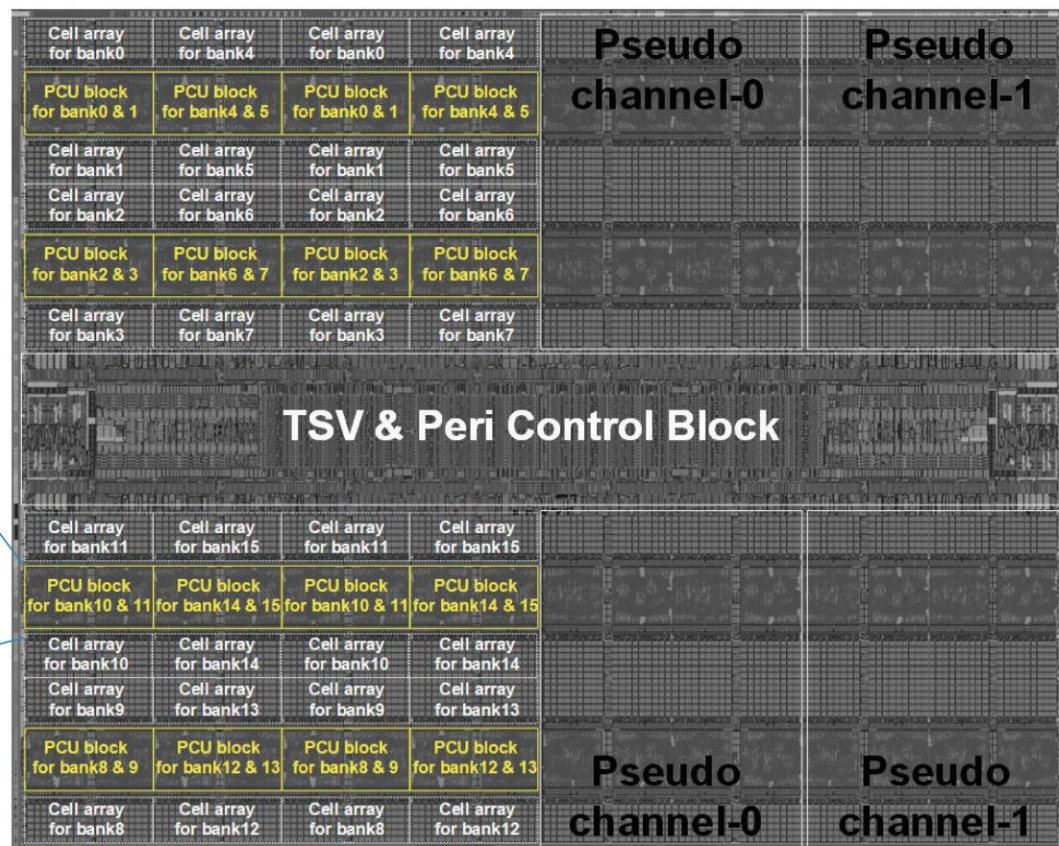
Samsung Function-in-Memory DRAM (2021)

Chip Implementation

- Mixed design methodology to implement FIMDRAM
 - Full-custom + Digital RTL

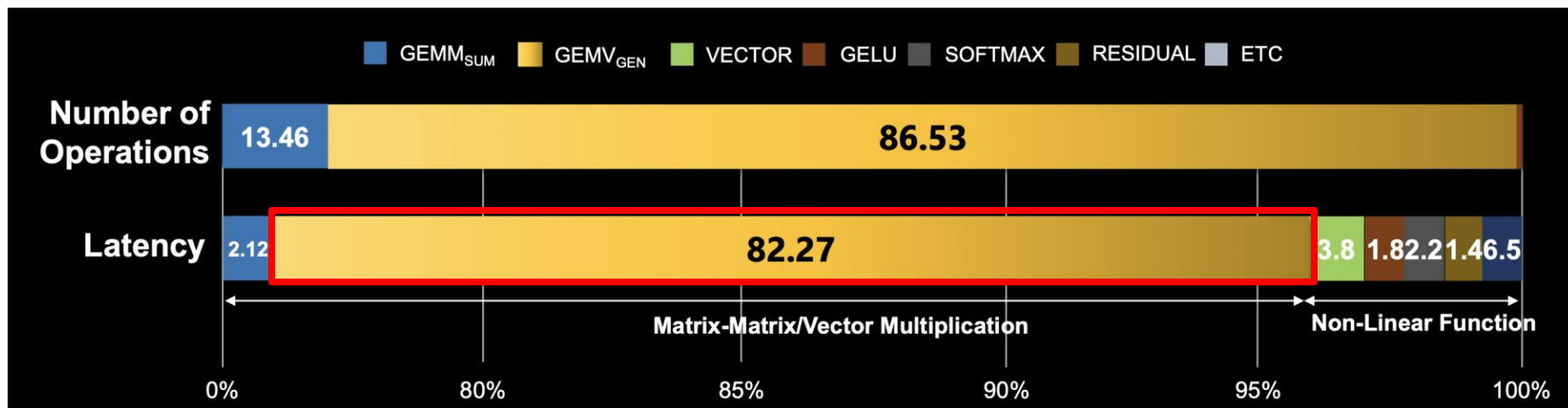


[Digital RTL design for PCU block]



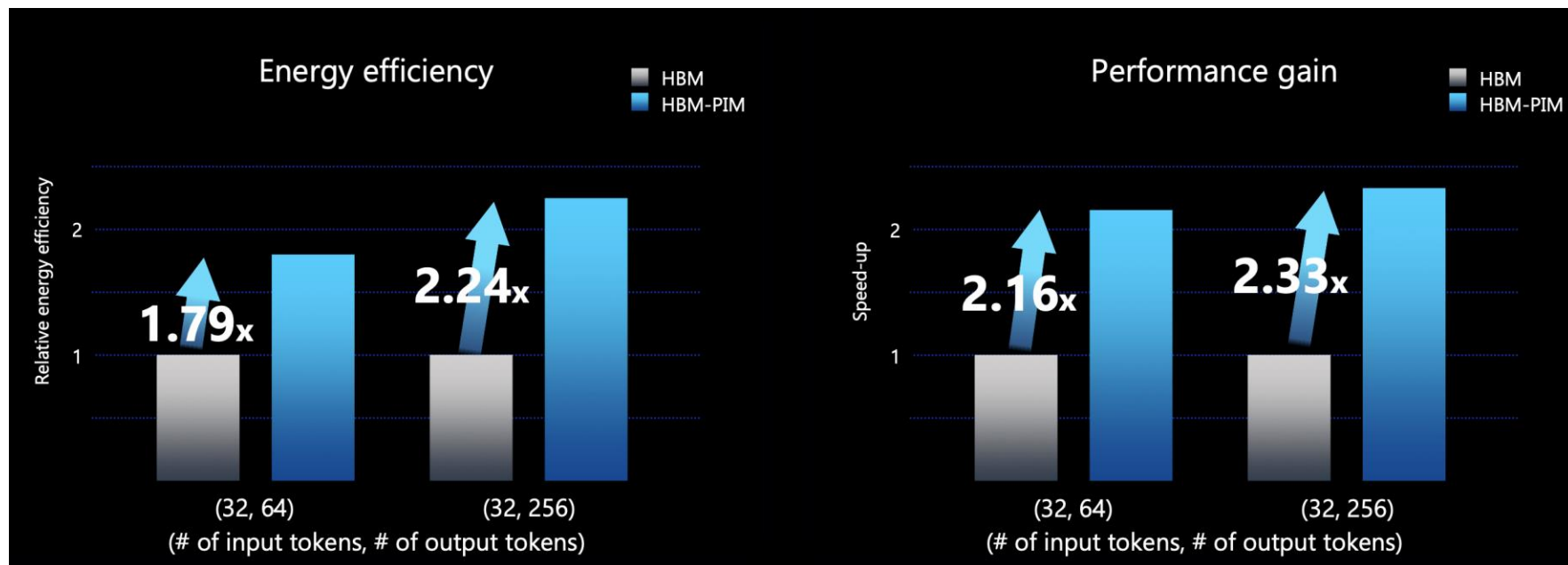
Samsung PNM Solutions for Generative AI (2023)

- Main target: **transformer** decoders used in **ChatGPT, GPT-3**
 - **Compute-bound step**: Summarization
 - **Memory-bound step**: Generation
 - Most of the execution time is spent on the **memory copy** from the **host CPU memory** to the **CPU memory**
- **GEMV** portion can be **60%-80% of total generation latency**, which is the target of PIM/PNM



Solution I: Samsung's HBM-PIM (2023)

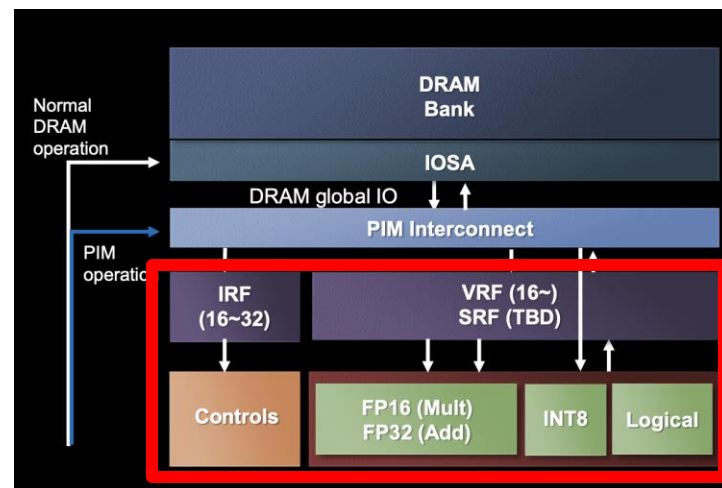
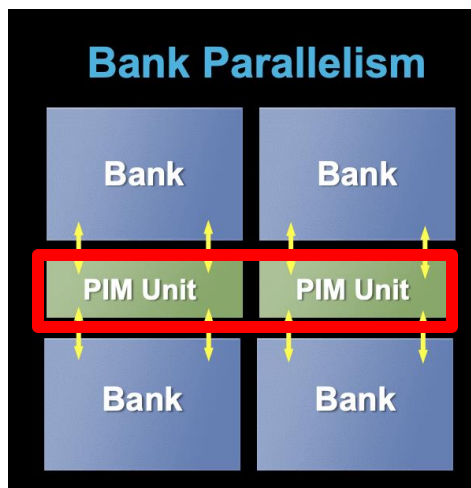
- AMD MI100 GPUs fabricated with HBM-PIM
- Experimental setup: GPT-J (6B, 32 input tokens), single AMD MI100-PIM GPU



- GPT can be accelerated by more than 2x over baseline

Solution II: Samsung's LPDDR-PIM (2023)

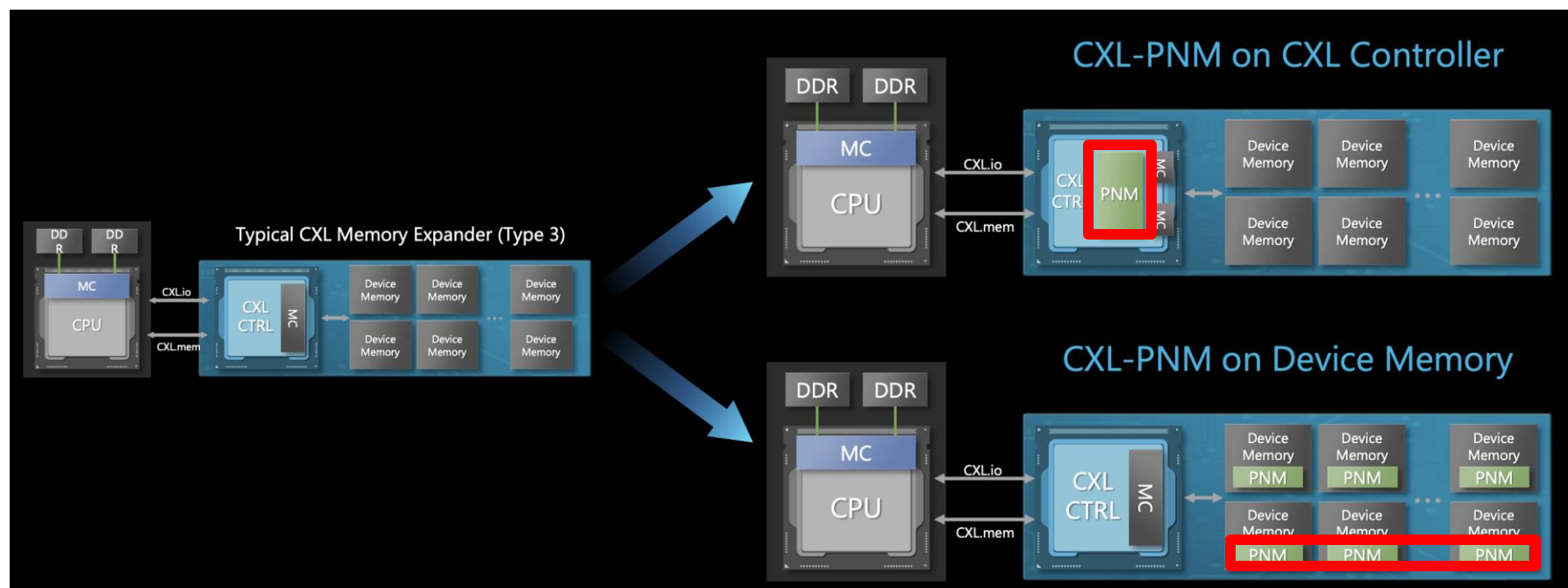
- PIM for on-device generative AI
 - ❑ Datacenter **costs** and **power consumption** are increasing due to the growing demand for cloud AI
- LPDDR-PIM improves **battery life** by preventing memory over-provisioning just for bandwidth



- 4.47x **performance gains** and 70.6% **energy reduction** in GPT-2

Solution III: Samsung's CXL-PNM (2023)

- A CXL-based processing-near-memory solution
 - ❑ Improves capacity, bandwidth, and power
 - ❑ Large-scale large-language models are often capacity-bound

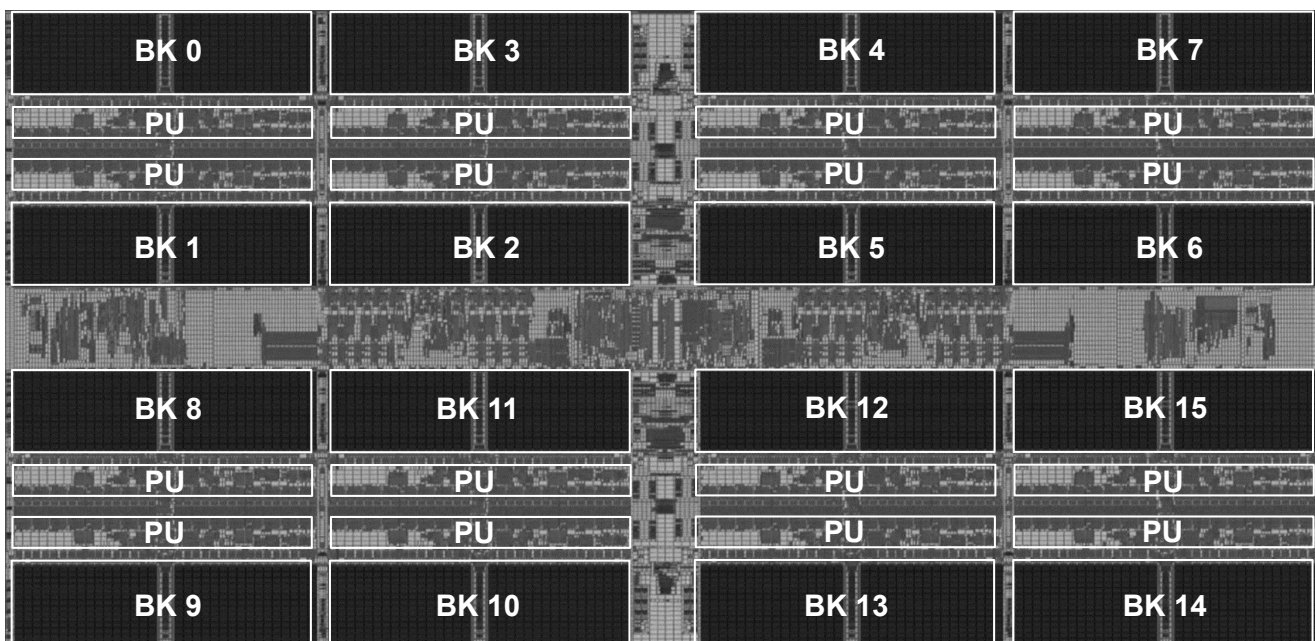


- Multiple CXL-PNM can offer 4.4x higher energy efficiency and 53% higher throughput than multiple GPUs

SK hynix AiM: Chip Implementation (2022)

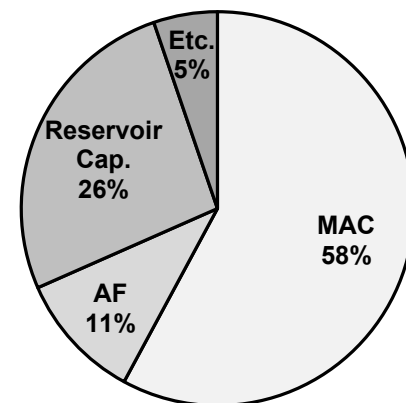
- 4 Gb AiM die with 16 processing units (PUs)

AiM Die Photograph



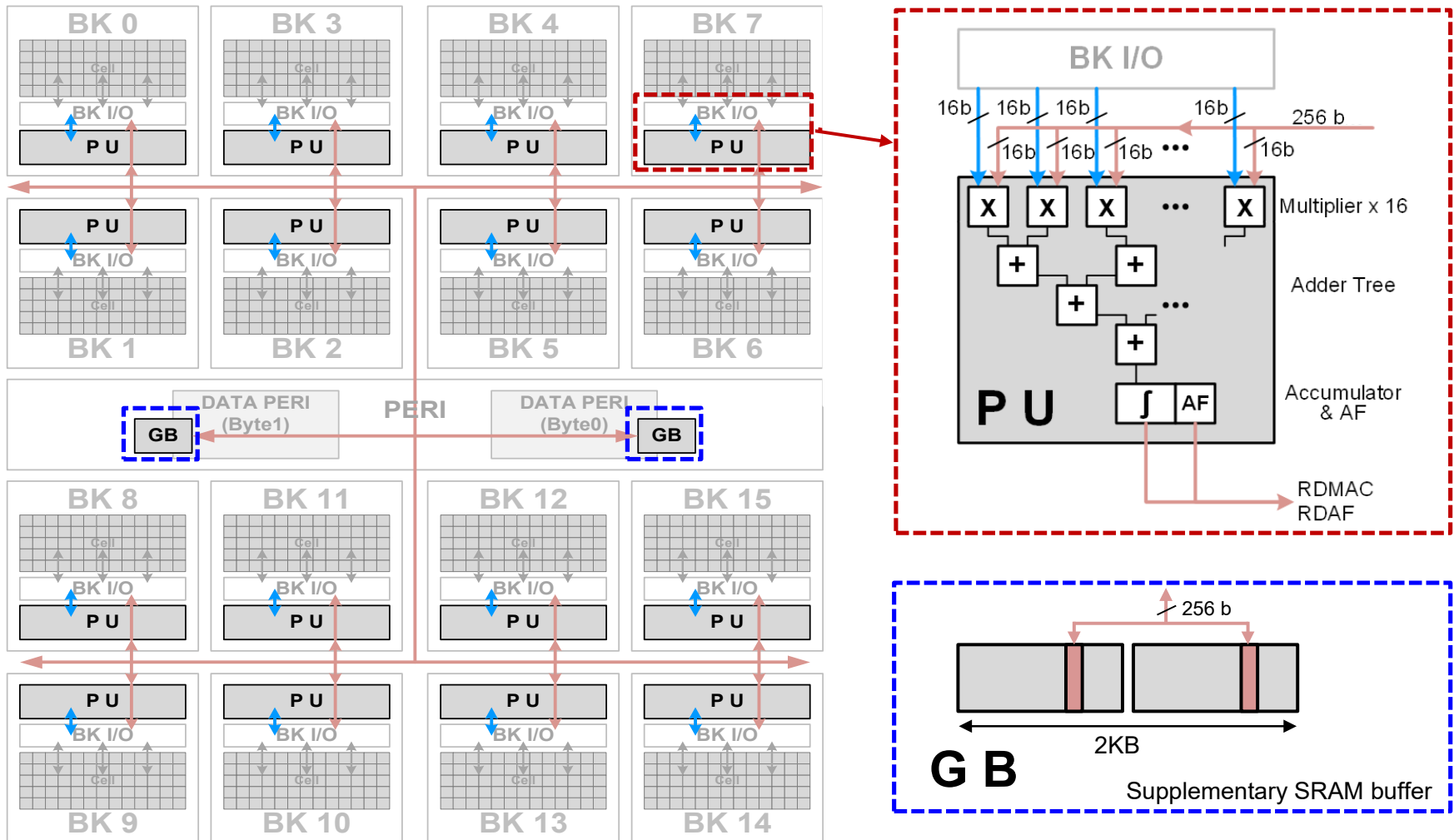
1 Process Unit (PU) Area

Total	0.19mm ²
MAC	0.11mm ²
Activation Function (AF)	0.02mm ²
Reservoir Cap.	0.05mm ²
Etc.	0.01mm ²



SK hynix AiM: System Organization (2022)

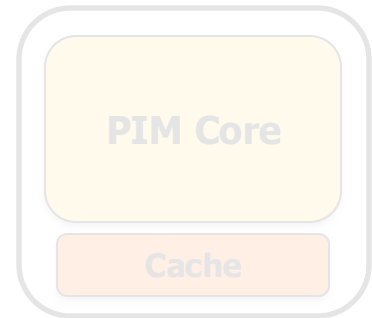
■ GDDR6-based AiM architecture



Possible PNM Designs

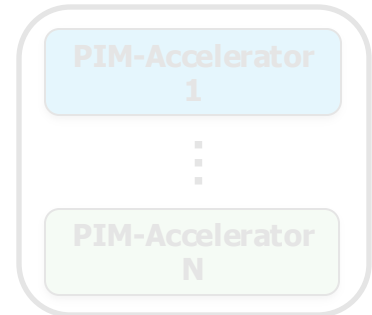
■ General-purpose programmable cores

- ❑ Wimpy cores (possibility of running any workload)
- ❑ E.g. from academia: Tesseract PIM for Graph Processing
- ❑ E.g. from industry: UPMEM PIM



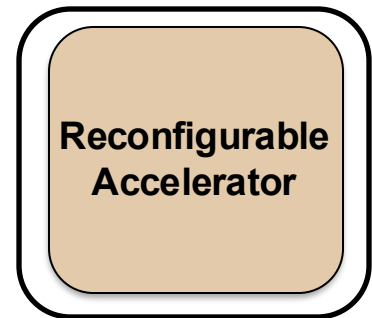
■ Fixed-function units

- ❑ Hardware/software co-designed PIM for efficiency
- ❑ E.g. from academia: Mensa for NN Edge Inference
- ❑ E.g. from industry: Samsung HBM-PIM, SK hynix AiM



■ Reconfigurable architectures

- ❑ PNM cores coupled with FPGAs, CGRA
- ❑ E.g. from academia: NERO for Weather Prediction
- ❑ E.g. from industry: Samsung AxDIMM



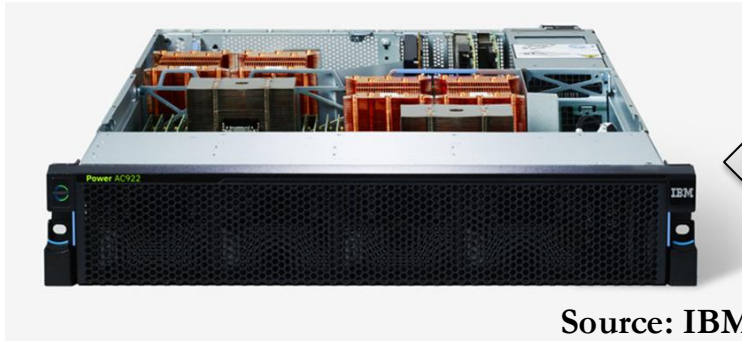
FPGA-based Processing Near Memory

- Gagandeep Singh, Dionysios Diamantopoulos, Christoph Hagleitner, Juan Gómez-Luna, Sander Stuijk, Onur Mutlu, and Henk Corporaal,
"NERO: A Near High-Bandwidth Memory Stencil Accelerator for Weather Prediction Modeling"
Proceedings of the 30th International Conference on Field-Programmable Logic and Applications (FPL), Gothenburg, Sweden, September 2020.

NERO: A Near High-Bandwidth Memory Stencil Accelerator for Weather Prediction Modeling

Gagandeep Singh^{a,b,c} Dionysios Diamantopoulos^c Christoph Hagleitner^c Juan Gómez-Luna^b
Sander Stuijk^a Onur Mutlu^b Henk Corporaal^a
^aEindhoven University of Technology ^bETH Zürich ^cIBM Research Europe, Zurich

Heterogeneous System: CPU+FPGA



Source: IBM

POWER9 AC922



Source: AlphaData

**DDR4-based AD9V3
board**

We evaluate two POWER9+FPGA systems:

1. HBM-based board AD9H7

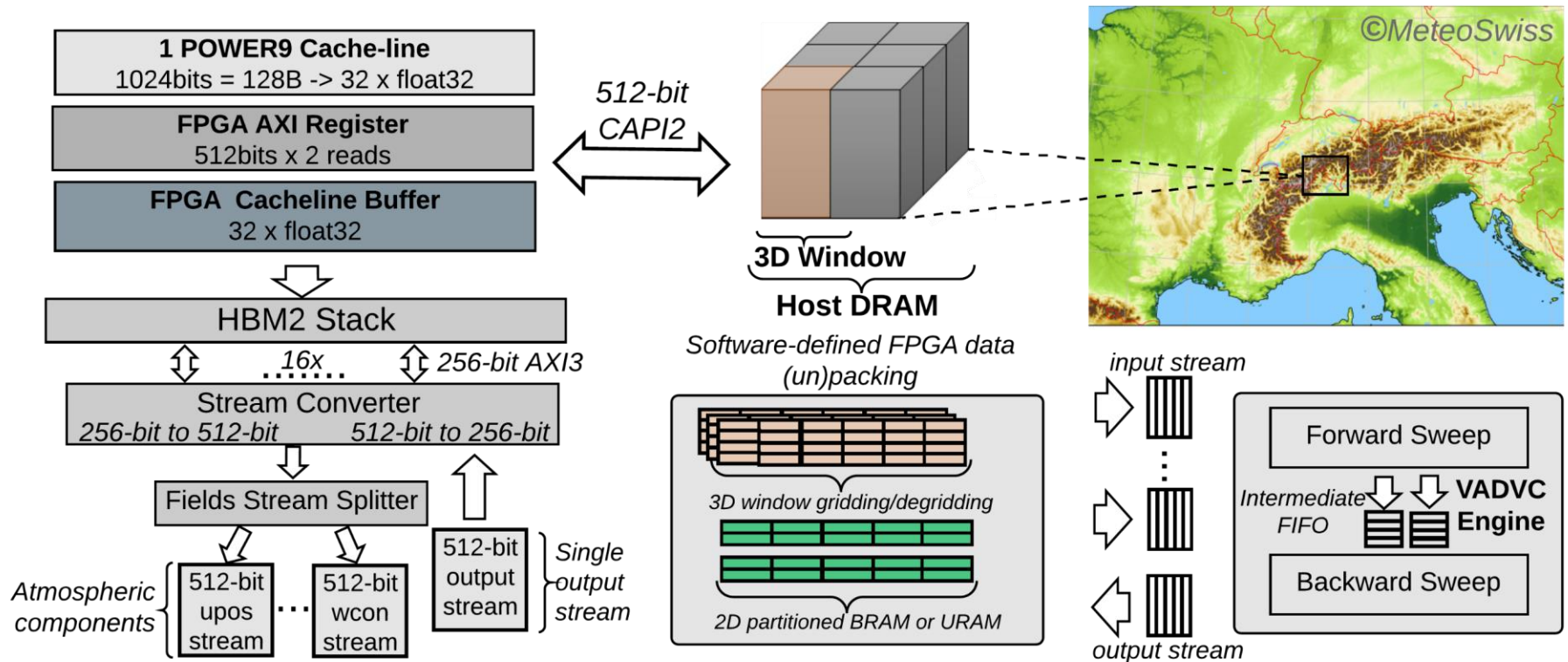
Xilinx Virtex Ultrascale+™ XCVU37P-2

2

2. DDR4-based board

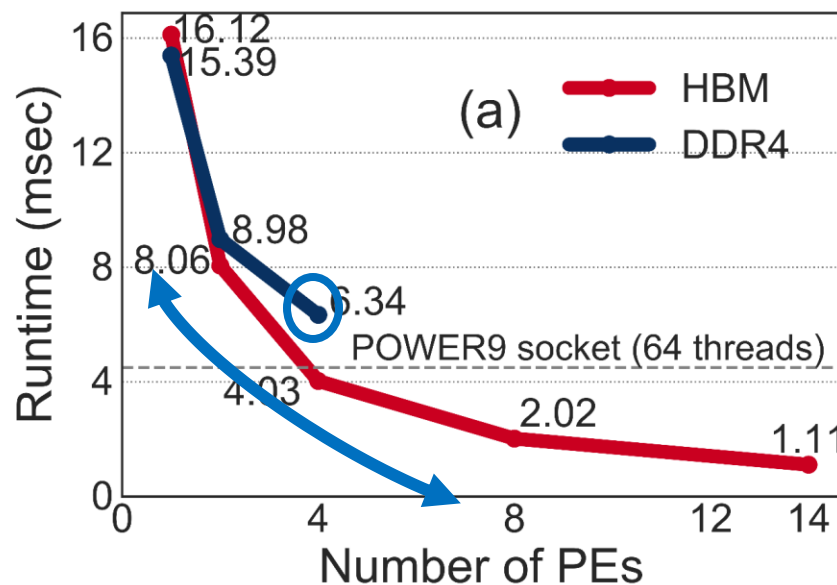
Xilinx Virtex Ultrascale+™ XCVU3P-

NERO Design Flow

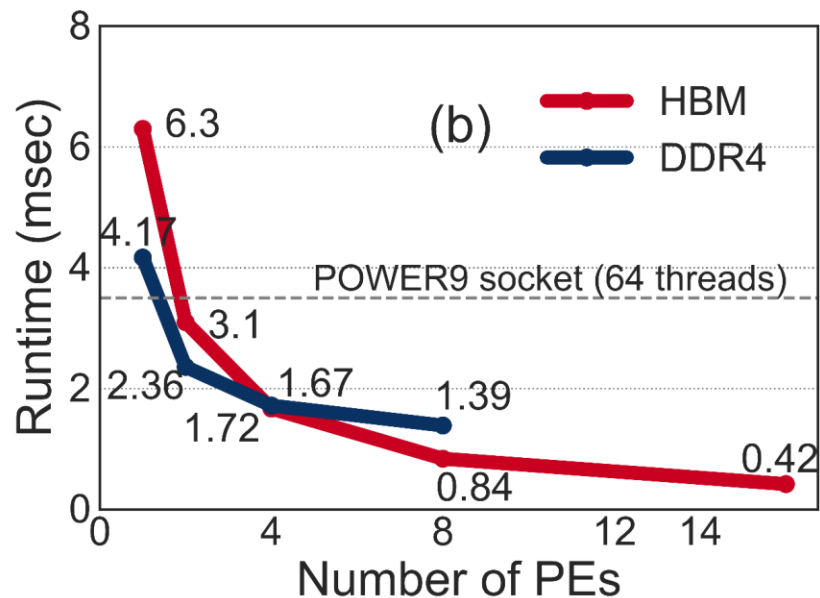


NERO Performance Analysis

Vertical Advection



Horizontal Diffusion

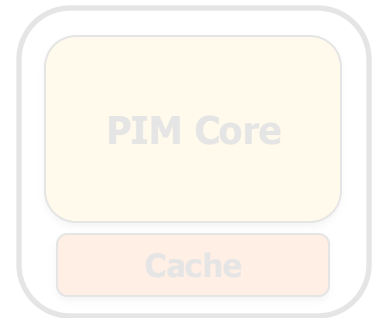


**NERO is 4.2x and 8.3x faster than
a complete POWER9 socket**

Possible PNM Designs

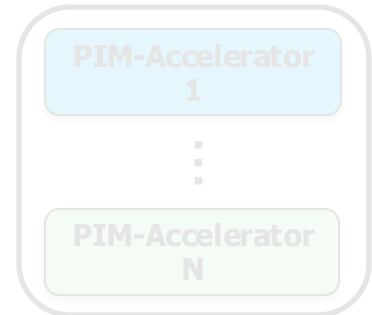
General-purpose programmable cores

- ❑ Wimpy cores (possibility of running any workload)
- ❑ E.g. from academia: Tesseract PIM for Graph Processing
- ❑ E.g. from industry: UPMEM PIM



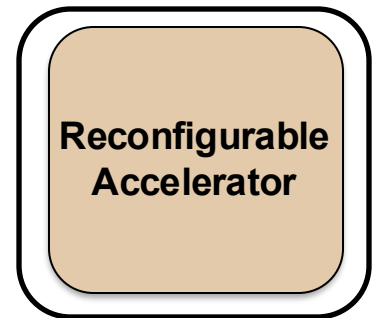
Fixed-function units

- ❑ Hardware/software co-designed PIM for efficiency
- ❑ E.g. from academia: Mensa for NN Edge Inference
- ❑ E.g. from industry: Samsung HBM-PIM, SK hynix AiM



Reconfigurable architectures

- ❑ PNM cores coupled with FPGAs, CGRA
- ❑ E.g. from academia: NERO for Weather Prediction
- ❑ E.g. from industry: Samsung AxDIMM



Samsung AxDIMM (2021)

Samsung Brings In-Memory Processing Power to Wider Range of Applications

Korea on August 24, 2021

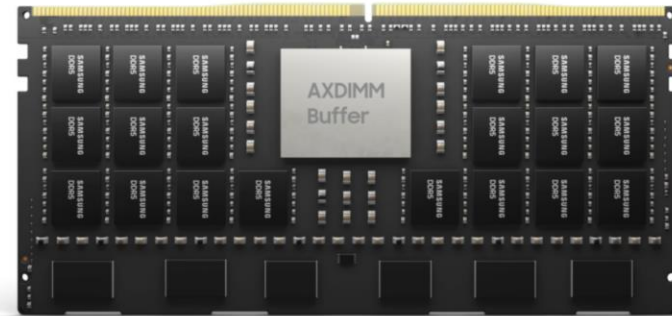
Audio Share

Integration of HBM-PIM with the Xilinx Alveo AI accelerator system will boost overall system performance by 2.5X while reducing energy consumption by more than 60%

PIM architecture will be broadly deployed beyond HBM, to include mainstream DRAM modules and mobile memory

Samsung Electronics, the world leader in advanced memory technology, today showcased its latest advancements with processing-in-memory (PIM) technology at Hot Chips 33—a leading semiconductor conference where the most notable microprocessor and IC innovations are unveiled each year. Samsung's revelations include the first successful integration of its PIM-enabled High Bandwidth Memory (HBM-PIM) into a commercialized accelerator system, and broadened PIM applications to embrace DRAM modules and mobile memory, in accelerating the move toward the convergence of memory and logic.

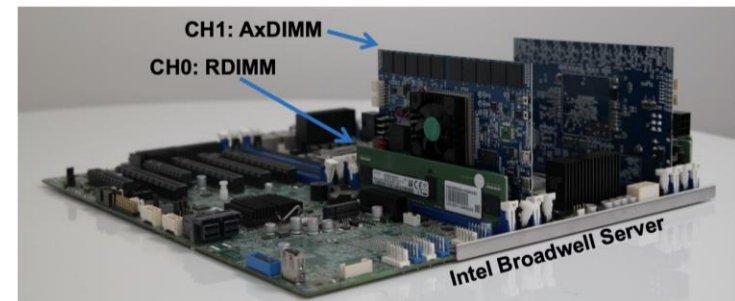
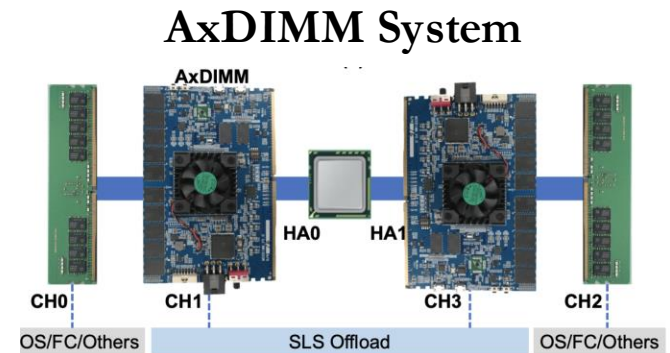
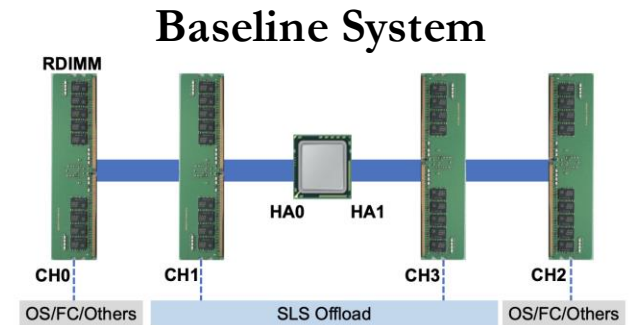
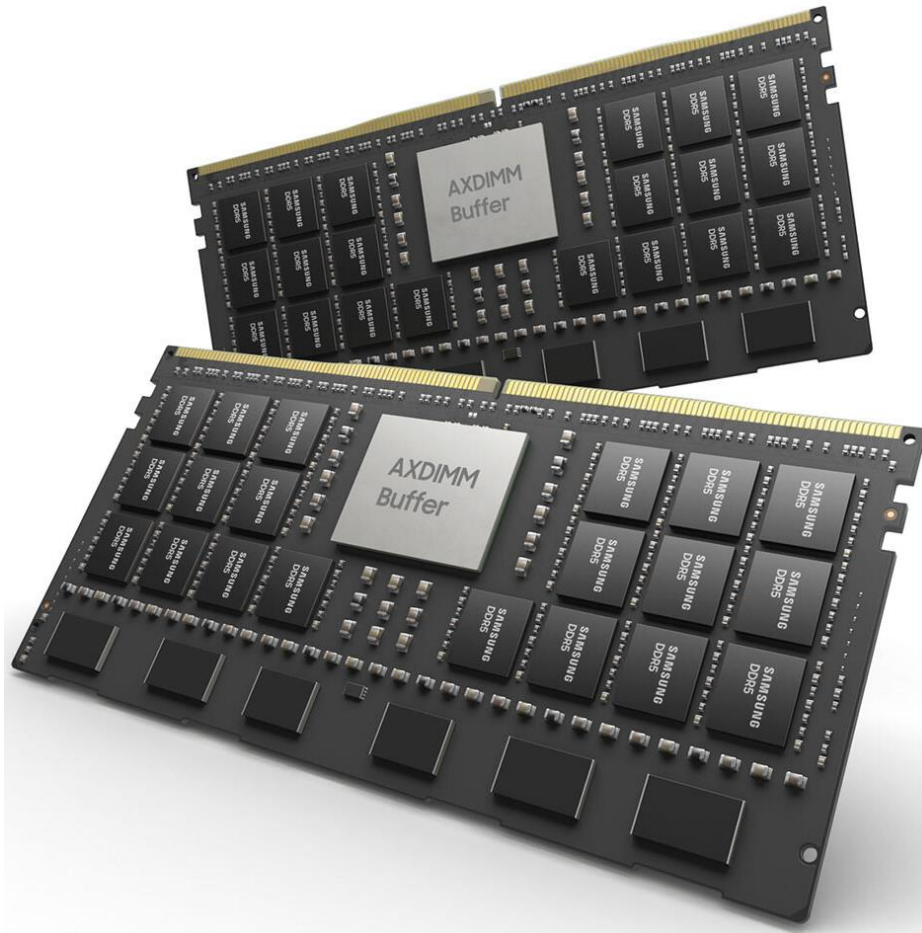
DRAM Modules Powered by PIM



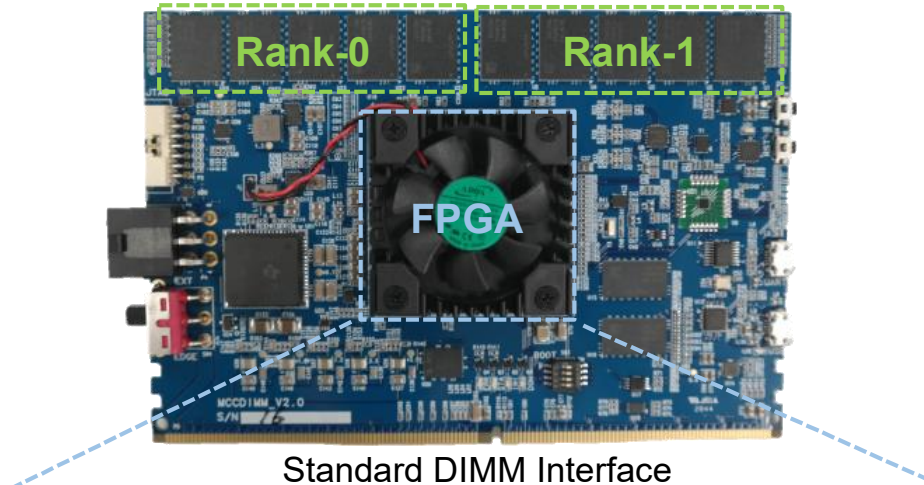
The Acceleration DIMM (AXDIMM) brings processing to the DRAM module itself, minimizing large data movement between the CPU and DRAM to boost the energy efficiency of AI accelerator systems. With an AI engine built inside the buffer chip, the AXDIMM can perform parallel processing of multiple memory ranks (sets of DRAM chips) instead of accessing just one rank at a time, greatly enhancing system performance and efficiency. Since the module can retain its traditional DIMM form factor, the AXDIMM facilitates drop-in replacement without requiring system modifications. Currently being tested on customer servers, the AXDIMM can offer approximately twice the performance in AI-based recommendation applications and a 40% decrease in system-wide energy usage.

Samsung AxDIMM (2021)

- DIMM-based PIM
 - DLRM recommendation system

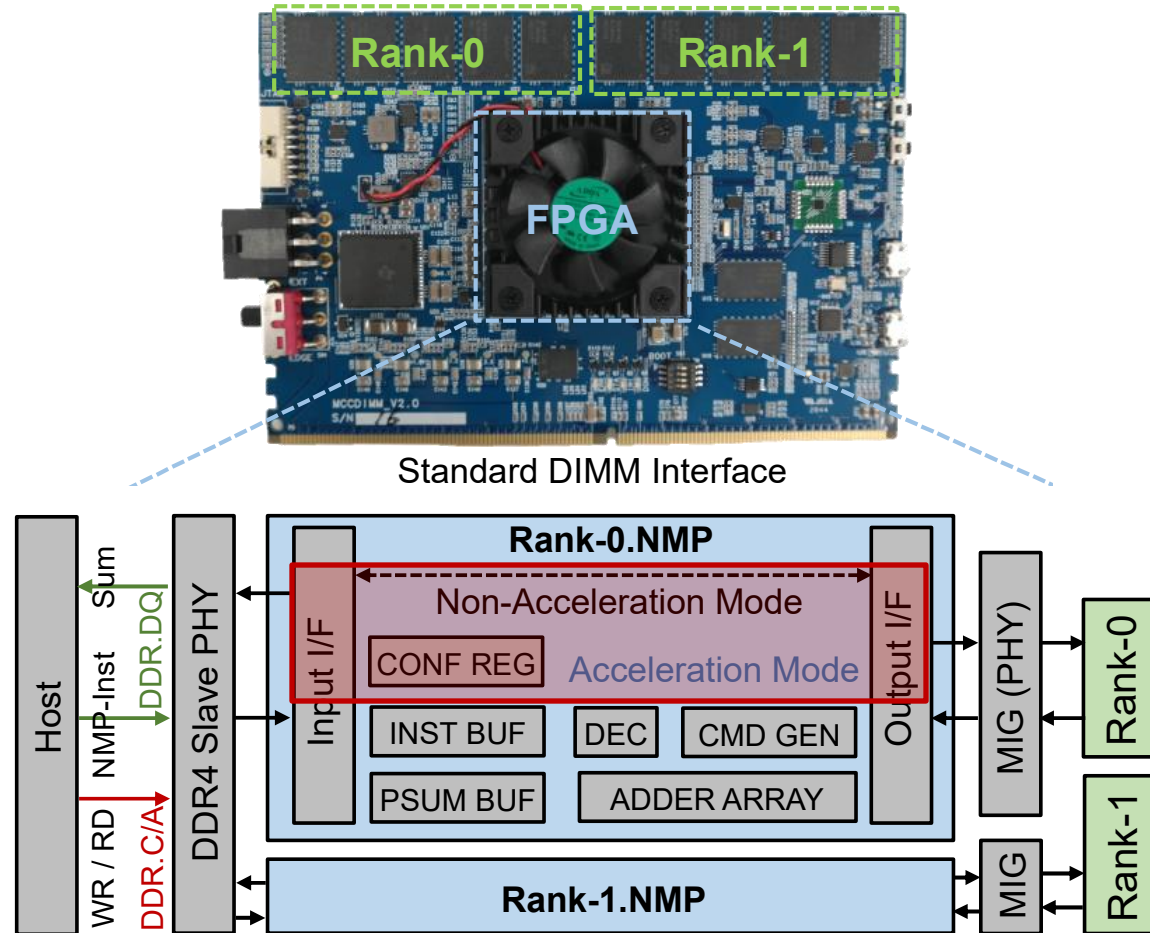


A_xDIMM Design: Hardware Architecture



FPGA board with standard DIMM interface:
It serves as a real-system
near-memory processing implementation

AxDIMM Design: Hardware Architecture



Two **execution modes**:
(1) non-acceleration mode
(2) acceleration mode (blocking)

Near-Memory Processing in Action: Accelerating Personalized Recommendation with AxDIMM

Liu Ke^{*†}, Xuan Zhang[†], Jinin So[‡], Jong-Geon Lee[‡], Shin-Haeng Kang[‡], Sukhan Lee[‡], Songyi Han[‡], YeonGon Cho[‡],
JIN Hyun Kim[‡], Yongsuk Kwon[‡], KyungSoo Kim[‡], Jin Jung[‡], Ilkwon Yun[‡], Sung Joo Park[‡], Hyunsun Park[‡],
Joonho Song[‡], Jeonghyeon Cho[‡], Kyomin Sohn[‡], Nam Sung Kim[‡], Hsien-Hsin S. Lee^{*}

^{*}Facebook, [†]Washington University in St. Louis, [‡]Samsung

An Architecture of Sparse Length Sum Accelerator in AxDIMM

Shin-haeng Kang
DRAM Design Team 1
Samsung Electronics
Hwasung, South Korea
s-h.kang@samsung.com

Byeongho Kim
DRAM Design Team 1
Samsung Electronics
Hwasung, South Korea
bh1122.kim@samsung.com

Sukhan Lee
DRAM Design Team 1
Samsung Electronics
Hwasung, South Korea
sh1026.lee@samsung.com

Kyomin Sohn
DRAM Design Team 1
Samsung Electronics
Hwasung, South Korea
kyomin.sohn@samsung.com

Improving In-Memory Database Operations with Acceleration DIMM (AxDIMM)

Donghun Lee
Minseon Ahn
Jungmin Kim
dong.hun.lee@sap.com
minseon.ahn@sap.com
jungmin.kim@sap.com
SAP Labs Korea

Jinin So
Jong-Geon Lee
Jeonghyeon Cho
Vishnu Charan Thummala
jinin.so@samsung.com
jg1021.lee@samsung.com
caleb1@samsung.com
vishnu.c.t@samsung.com
Samsung Electronics

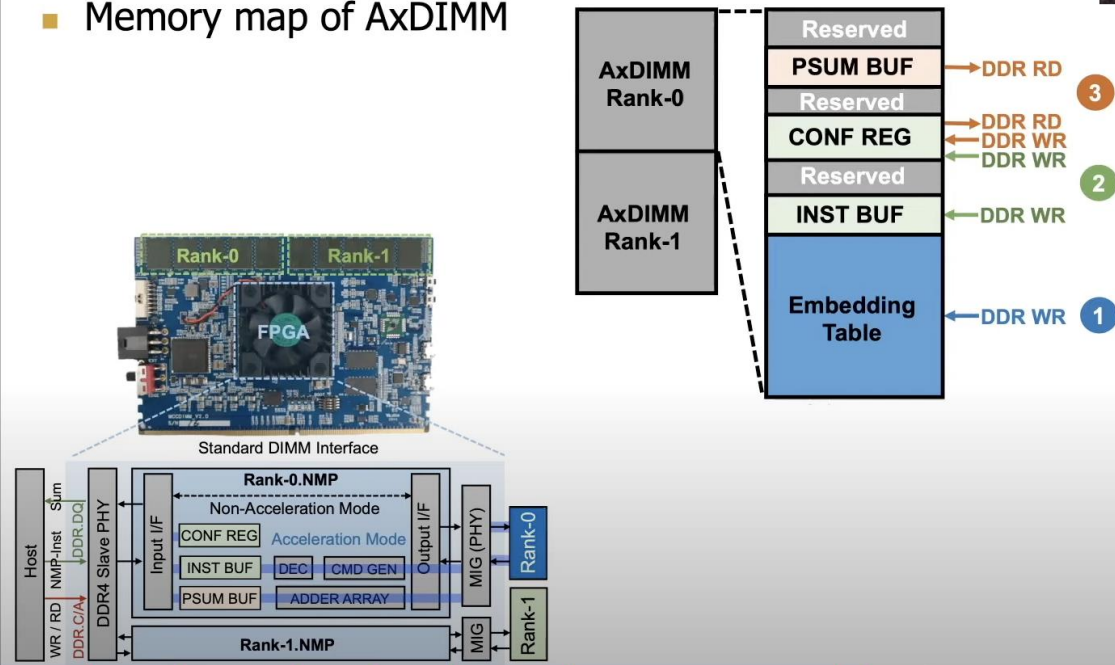
Oliver Rebholz
oliver.rebholz@sap.com
SAP SE

Ravi Shankar JV
Sachin Suresh Upadhyaya
Mohammed Ibrahim Khan
Jin Hyun Kim
venkata.ravi@samsung.com
sachin1.s@samsung.com
ibrahim.khan@samsung.com
kjh5555@samsung.com
Samsung Electronics

Longer Lecture on AxDIMM

AxDIMM Design: Address Map

■ Memory map of AxDIMM



PIM Course: Lecture 7: Real-world PIM: Samsung AxDIMM - Fall 2022



Onur Mutlu Lectures

32.4K subscribers



Subscribed

21



Share

Clip

Save



846 views 4 months ago Livestream - P&S Data-Centric Architectures: Fundamentally Improving Performance and Energy (Fall 2022)

Projects & Seminars, ETH Zürich, Fall 2022

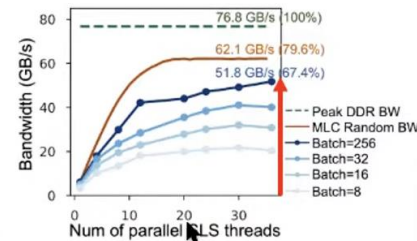
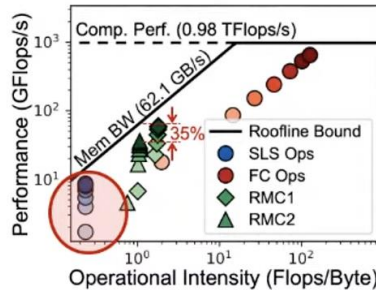
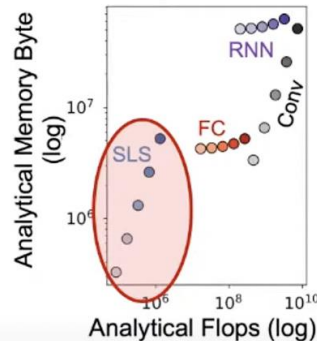
Data-Centric Architectures: Fundamentally Improving Performance and Energy

(https://safari.ethz.ch/projects_and_s...) Show more

Another Longer Lecture on AxDIMM

DLRM Performance Characterization

- Identifying **key performance bottlenecks** for the DLRM system



SparseLengths (SLS) operators:

- Low FP intensity
- Larger batch size:
 - Higher memory footprint
 - Higher memory intensity

The **memory bandwidth can easily be saturated** by embedding operations especially as both the batch size and the number of threads increase

Processing in Memory Course: Meeting 5: Real-world PIM architectures IV - Fall'21



Onur Mutlu Lectures

32.4K subscribers



Subscribed

18



Share

Clip

Save



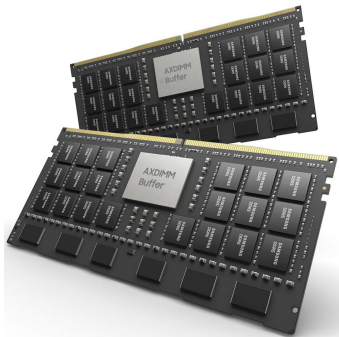
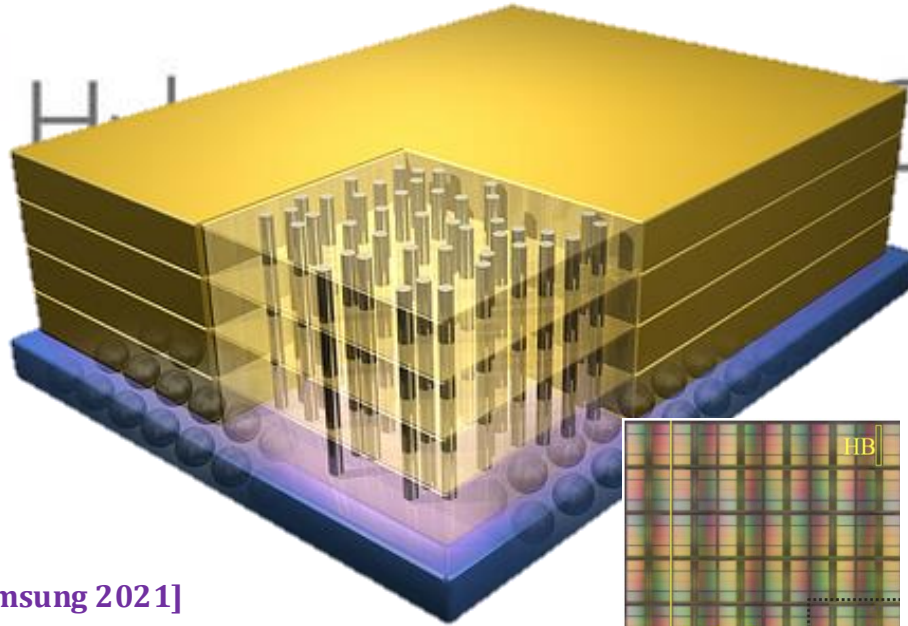
676 views Streamed 1 year ago Livestream - P&S Exploring the Processing-in-Memory Paradigm for Future Computing Systems (Fall 2021)

Project & Seminar, ETH Zürich, Fall 2021

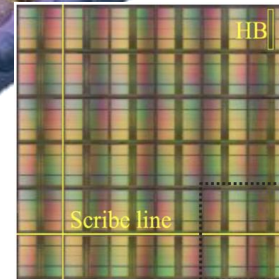
Exploring the Processing-in-Memory Paradigm for Future Computing Systems (https://safari.ethz.ch/projects_and_s...)

Show more

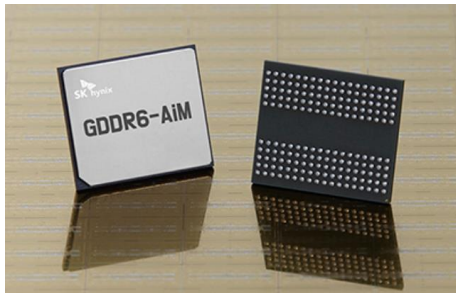
Processing-in-Memory Landscape Today



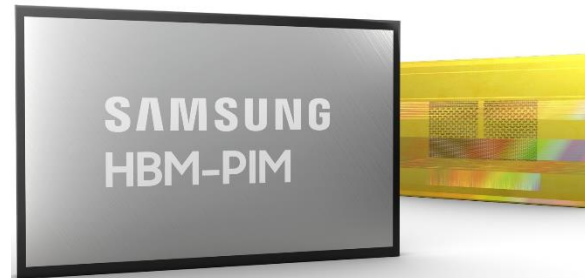
[Samsung 2021]



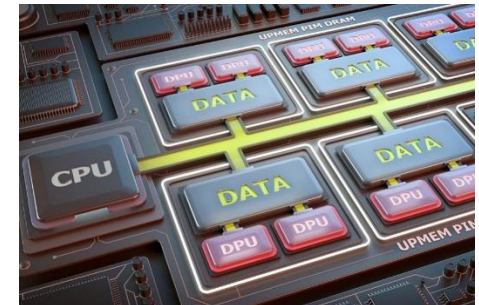
[Alibaba 2022]



[SK Hynix 2022]



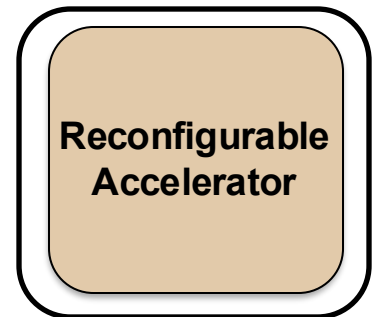
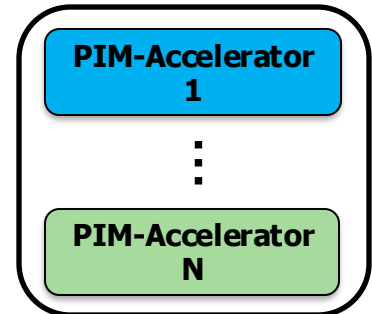
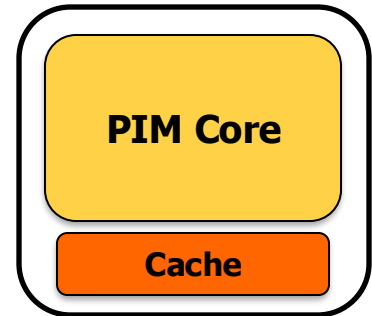
[Samsung 2021]



[UPMEM 2019]

Possible PNM Designs

- **General-purpose** programmable cores
 - ❑ Wimpy cores (possibility of running any workload)
 - ❑ E.g. from academia: Tesseract PIM for Graph Processing
 - ❑ E.g. from industry: UPMEM PIM
- **Fixed-function** units
 - ❑ Hardware/software co-designed PIM for efficiency
 - ❑ E.g. from academia: Mensa for NN Edge Inference
 - ❑ E.g. from industry: Samsung HBM-PIM, SK hynix AiM
- **Reconfigurable** architectures
 - ❑ PNM cores coupled with FPGAs, CGRA
 - ❑ E.g. from academia: NERO for Weather Prediction
 - ❑ E.g. from industry: Samsung AxDIMM



Research Tools PNM: DAMOV-SIM

- Geraldo F. Oliveira, Juan Gomez-Luna, Lois Orosa, Saugata Ghose, Nandita Vijaykumar, Ivan fernandez, Mohammad Sadrosadati, and Onur Mutlu,
"DAMOV: A New Methodology and Benchmark Suite for Evaluating Data Movement Bottlenecks"
IEEE Access, 8 September 2021.
*Preprint in **arXiv***, 8 May 2021.
[[arXiv preprint](#)]
[[IEEE Access version](#)]
[[DAMOV Suite and Simulator Source Code](#)]
[[SAFARI Live Seminar Video](#) (2 hrs 40 mins)]
[[Short Talk Video](#) (21 minutes)]

DAMOV: A New Methodology and Benchmark Suite for Evaluating Data Movement Bottlenecks

GERALDO F. OLIVEIRA, ETH Zürich, Switzerland

JUAN GÓMEZ-LUNA, ETH Zürich, Switzerland

LOIS OROSA, ETH Zürich, Switzerland

SAUGATA GHOSE, University of Illinois at Urbana–Champaign, USA

NANDITA VIJAYKUMAR, University of Toronto, Canada

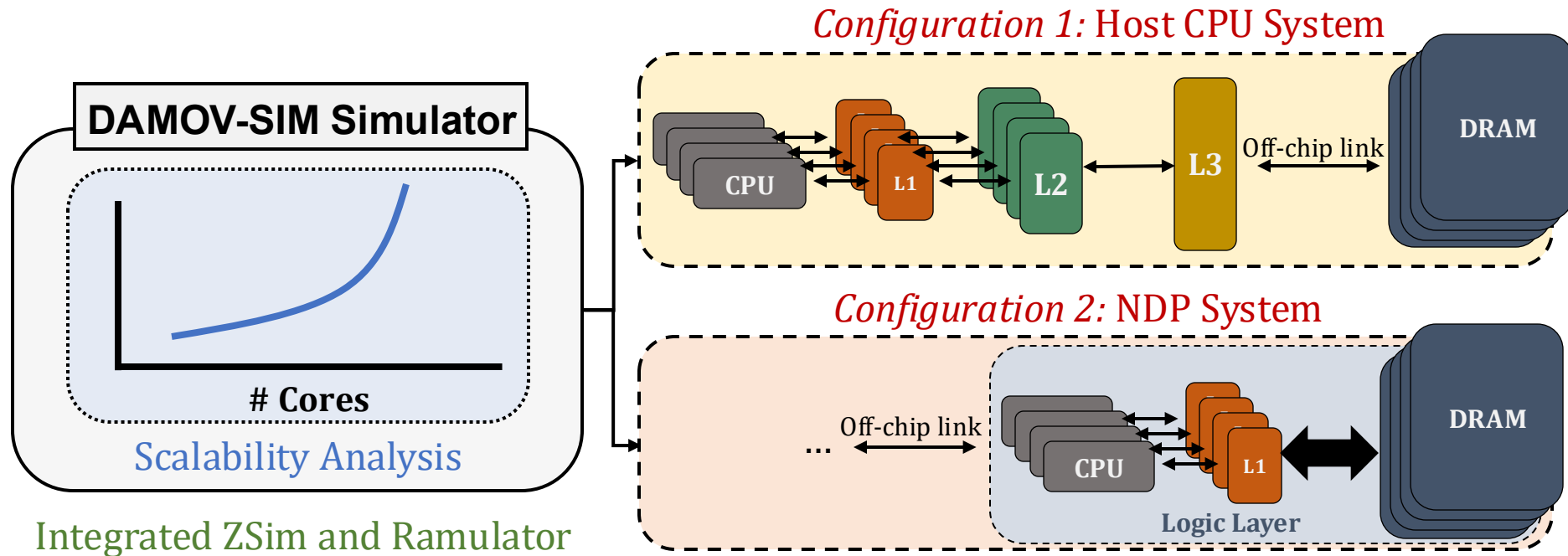
IVAN FERNANDEZ, University of Malaga, Spain & ETH Zürich, Switzerland

MOHAMMAD SADROSADATI, ETH Zürich, Switzerland

ONUR MUTLU, ETH Zürich, Switzerland

Step 3: Memory Bottleneck Classification (2/2)

- **Goal:** identify the specific sources of data movement bottlenecks



- **Scalability Analysis:**
 - 1, 4, 16, 64, and 256 out-of-order/in-order host and NDP CPU cores
 - 3D-stacked memory as main memory

DAMOV is Open Source

- We open-source our **benchmark suite** and our **toolchain**

CMU-SAFARI / DAMOV

<> Code Issues Pull requests Actions Projects Security Insights Settings

main 1 branch 0 tags

Go to file

Add file

Code

About



DAMOV is a benchmark suite and a methodical framework targeting the study of data movement bottlenecks in modern applications. It is intended to study new architectures, such as near-data processing. Described by Oliveira et al. (preliminary version at <https://arxiv.org/pdf/2105.03725.pdf>)

Readme

Releases

No releases published
[Create a new release](#)

Packages

No packages published
[Publish your first package](#)

Languages



omutlu Update README.md

ce1b4ea 17 days ago 5 commits

simulator	Cleaning	19 days ago
README.md	Update README.md	17 days ago
get_workloads.sh	DAMOV -- first commit	19 days ago

README.md

DAMOV: A New Methodology and Benchmark Suite for Evaluating Data Movement Bottlenecks

DAMOV is a benchmark suite and a methodical framework targeting the study of data movement bottlenecks in modern applications. It is intended to study new architectures, such as near-data processing.

The DAMOV benchmark suite is the first open-source benchmark suite for main memory data movement-related studies, based on our systematic characterization methodology. This suite consists of 144 functions representing different sources of data movement bottlenecks and can be used as a baseline benchmark set for future data-movement mitigation research. The applications in the DAMOV benchmark suite belong to popular benchmark suites, including [BWA](#), [Chai](#), [Darknet](#), [GASE](#), [Hardware Effects](#), [Hashjoin](#), [HPCC](#), [HPCG](#), [Ligra](#), [PARSEC](#), [Parboil](#), [PolyBench](#), [Phoenix](#), [Rodinia](#), [SPLASH-2](#), [STREAM](#).

DAMOV-SIM

DAMOV
Benchmarks

SAFARI

DAMOV is Open Source

- We open-source our [benchmark suite](https://github.com/CMU-SAFARI/DAMOV) and our [toolchain](https://github.com/CMU-SAFARI/DAMOV)

CMU-SAFARI / DAMOV

<> Code Issues Pull requests Actions Projects Security Insights Settings

main 1 branch 0 tags

Go to file

Add file

Code

About

DAMOV is a benchmark suite and a

Get DAMOV at:

<https://github.com/CMU-SAFARI/DAMOV>

README.md

DAMOV: A New Methodology and Benchmark Suite for Evaluating Data Movement Bottlenecks

DAMOV is a benchmark suite and a methodical framework targeting the study of data movement bottlenecks in modern applications. It is intended to study new architectures, such as near-data processing.

The DAMOV benchmark suite is the first open-source benchmark suite for main memory data movement-related studies, based on our systematic characterization methodology. This suite consists of 144 functions representing different sources of data movement bottlenecks and can be used as a baseline benchmark set for future data-movement mitigation research. The applications in the DAMOV benchmark suite belong to popular benchmark suites, including [BWA](#), [Chai](#), [Darknet](#), [GASE](#), [Hardware Effects](#), [Hashjoin](#), [HPCC](#), [HPCG](#), [Ligra](#), [PARSEC](#), [Parboil](#), [PolyBench](#), [Phoenix](#), [Rodinia](#), [SPLASH-2](#), [STREAM](#).

Readme

Releases

No releases published
[Create a new release](#)

Packages

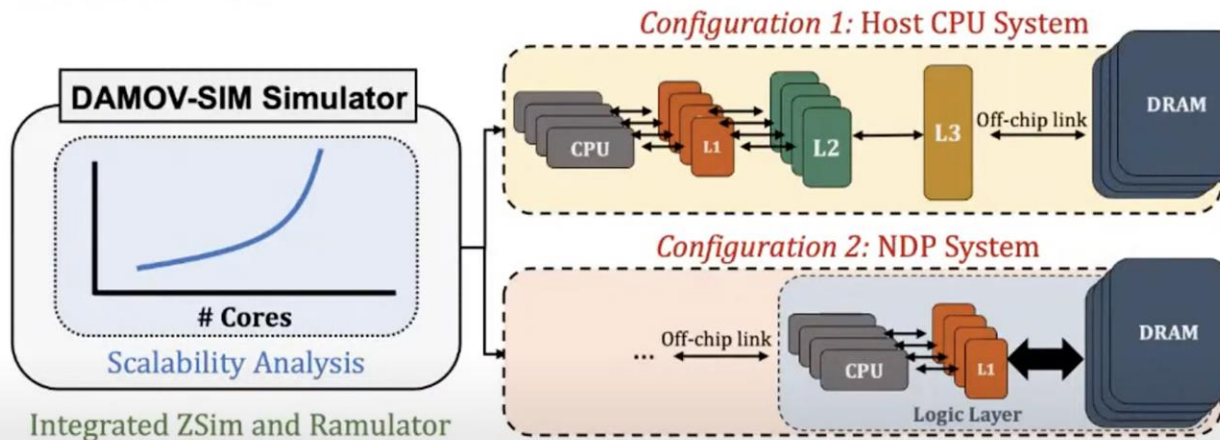
No packages published
[Publish your first package](#)

Languages

More on DAMOV Analysis Methodology & Workloads

Step 3: Memory Bottleneck Classification (2/)

- **Goal:** identify the specific sources of data movement bottlenecks



- **Scalability Analysis:**
 - 1, 4, 16, 64, and 256 out-of-order/in-order host and NDP CPU cores
 - 3D-stacked memory as main memory

DAMOV-SIM: <https://github.com/CMU-SAFARI/DAMOV>

SAFARI Live Seminar: DAMOV: A New Methodology & Benchmark Suite for Data Movement Bottlenecks

352 views • Streamed live on Jul 22, 2021

18 0 SHARE SAVE ...



Onur Mutlu Lectures
17.7K subscribers

ANALYTICS

EDIT VIDEO

More on DAMOV Methods & Benchmarks

- Geraldo F. Oliveira, Juan Gomez-Luna, Lois Orosa, Saugata Ghose, Nandita Vijaykumar, Ivan fernandez, Mohammad Sadrosadati, and Onur Mutlu,
["DAMOV: A New Methodology and Benchmark Suite for Evaluating Data Movement Bottlenecks"](#)
[IEEE Access](#), 8 September 2021.
Preprint in [arXiv](#), 8 May 2021.
[[arXiv preprint](#)]
[[IEEE Access version](#)]
[[DAMOV Suite and Simulator Source Code](#)]
[[SAFARI Live Seminar Video](#) (2 hrs 40 mins)]
[[Short Talk Video](#) (21 minutes)]

DAMOV: A New Methodology and Benchmark Suite for Evaluating Data Movement Bottlenecks

GERALDO F. OLIVEIRA, ETH Zürich, Switzerland

JUAN GÓMEZ-LUNA, ETH Zürich, Switzerland

LOIS OROSA, ETH Zürich, Switzerland

SAUGATA GHOSE, University of Illinois at Urbana–Champaign, USA

NANDITA VIJAYKUMAR, University of Toronto, Canada

IVAN FERNANDEZ, University of Malaga, Spain & ETH Zürich, Switzerland

MOHAMMAD SADROSADATI, ETH Zürich, Switzerland

ONUR MUTLU, ETH Zürich, Switzerland

Research Tools PNM: Samsung HBM-PIM

<https://github.com/SAITPublic/PIMSimulator>

SAITPublic / PIMSimulator Public

Notifications F

<> Code Issues 14 Pull requests 2 Actions Projects Security Insights

dev 1 Branch Tags

Go to file

<> Code

iamshcha Reformat and some fixes 8a90a6e · 5 months ago 14 Commits

data	Initial commit	2 years ago
dump	Initial commit	2 years ago
ini	Initial commit	2 years ago
lib	Initial commit	2 years ago
src	Reformat and some fixes	5 months ago
tools/emulator_api	Modified to swap rows only for emulation paths	last year
.gitignore	Initial commit	2 years ago
LICENSE-DRAMSIM2	Initial commit	2 years ago
LICENSE-PIMSimulator	Initial commit	2 years ago
README.md	Modified the prerequisite installation method	9 months ago
Sconstruct	Initial commit	2 years ago
system_hbm.ini	Initial commit	2 years ago
system_hbm_1ch.ini	Initial commit	2 years ago
system_hbm_64ch.ini	Initial commit	2 years ago

About

Processing-In-Memory (PIM) Simulator

Readme

View license

Activity

Custom properties

126 stars

3 watching

42 forks

Report repository

Releases

No releases published

Packages

No packages published

Languages

C++ 99.6%

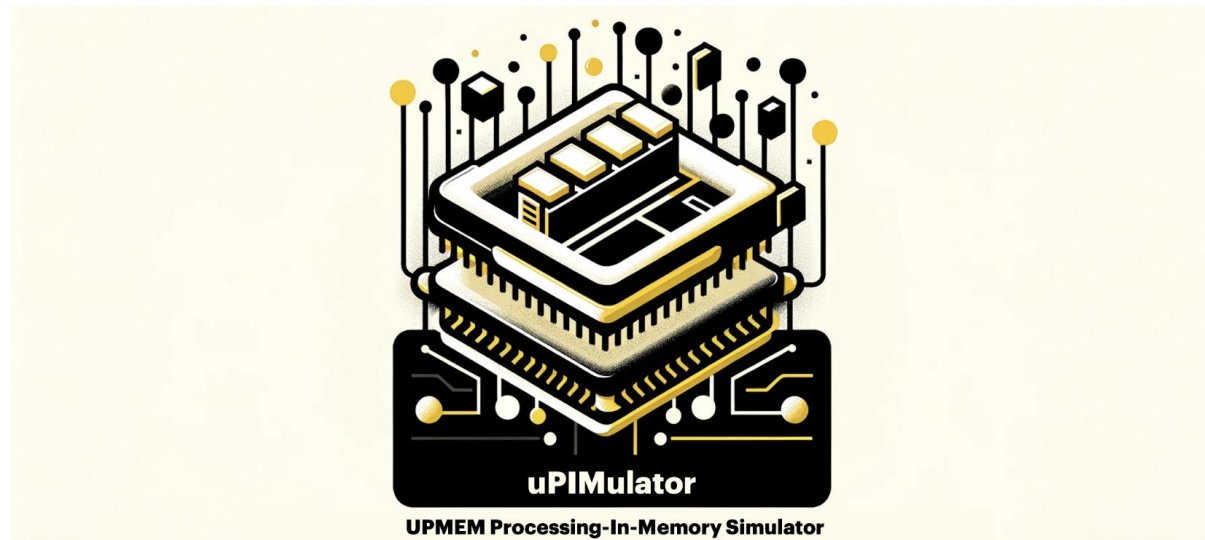
Python 0.4%

Research Tools PNM: UPMEM PIM (I)

<https://github.com/VIA-Research/uPIMulator>

📖 README 📄 MIT license

📖 Introduction



Welcome to the uPIMulator Framework Documentation!

This documentation serves as your comprehensive guide to the uPIMulator framework, catering to both novice and experienced researchers. Here, you'll find the resources necessary to leverage uPIMulator effectively for your research projects.

We provide in-depth coverage of uPIMulator's features, from foundational concepts to advanced functionalities. Explore this documentation to unlock the full potential of uPIMulator and elevate your research endeavors.

Research Tools PNM: UPMEM PIM (II)

<https://ieeexplore.ieee.org/document/10476411/>

2024 IEEE International Symposium on High-Performance Computer Architecture (HPCA)

Pathfinding Future PIM Architectures by Demystifying a Commercial PIM Technology

Bongjoon Hyun Taehun Kim Dongjae Lee Minsoo Rhu

KAIST

`{bongjoon.hyun, taehun.kim, dongjae.lee, mrhu}@kaist.ac.kr`

2nd Workshop on Memory-Centric Computing: Processing-Near-Memory - Part I

Dr. Geraldo F. Oliveira

<https://geraldofojunior.github.io>

ICS 2025

08 June 2025