

Leveraging the PULP Platform to build Reliable Multicore SoCs for Space Applications

Michael Rogenmoser

Yvan Tortorella

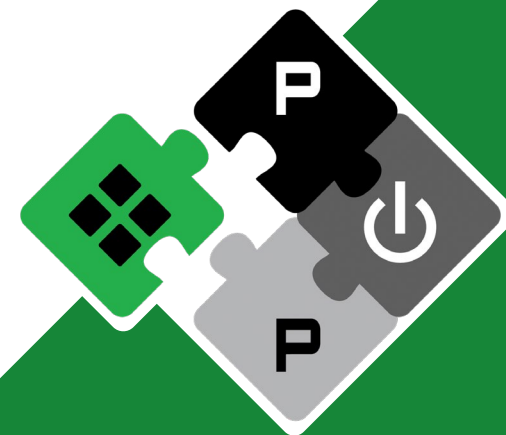
and the PULP team

michaero@iis.ee.ethz.ch

yvan.tortorella@unibo.it

PULP Platform

Open Source Hardware, the way it should be!



@pulp_platform 

pulp-platform.org 

youtube.com/pulp_platform 

Overview



1. The PULP Platform
2. Leveraging the PULP cluster for a Fault-Tolerant System
3. Using Heterogeneity for High-Performance Systems

Who are we and who is behind PULP?



Michael



Yvan



ALMA MATER STUDIORUM
UNIVERSITA DI BOLOGNA

ETH zürich



Dr. Frank Gurkaynak Prof. Luca Benini Prof. Davide Rossi



ETH zürich



ALMA MATER STUDIORUM
UNIVERSITA DI BOLOGNA

Michael Rogenmoser, Yvan Tortorella - ESA RISC-V in Space



PULP: Energy-Efficient Computing from Core to Platform



- **Core**

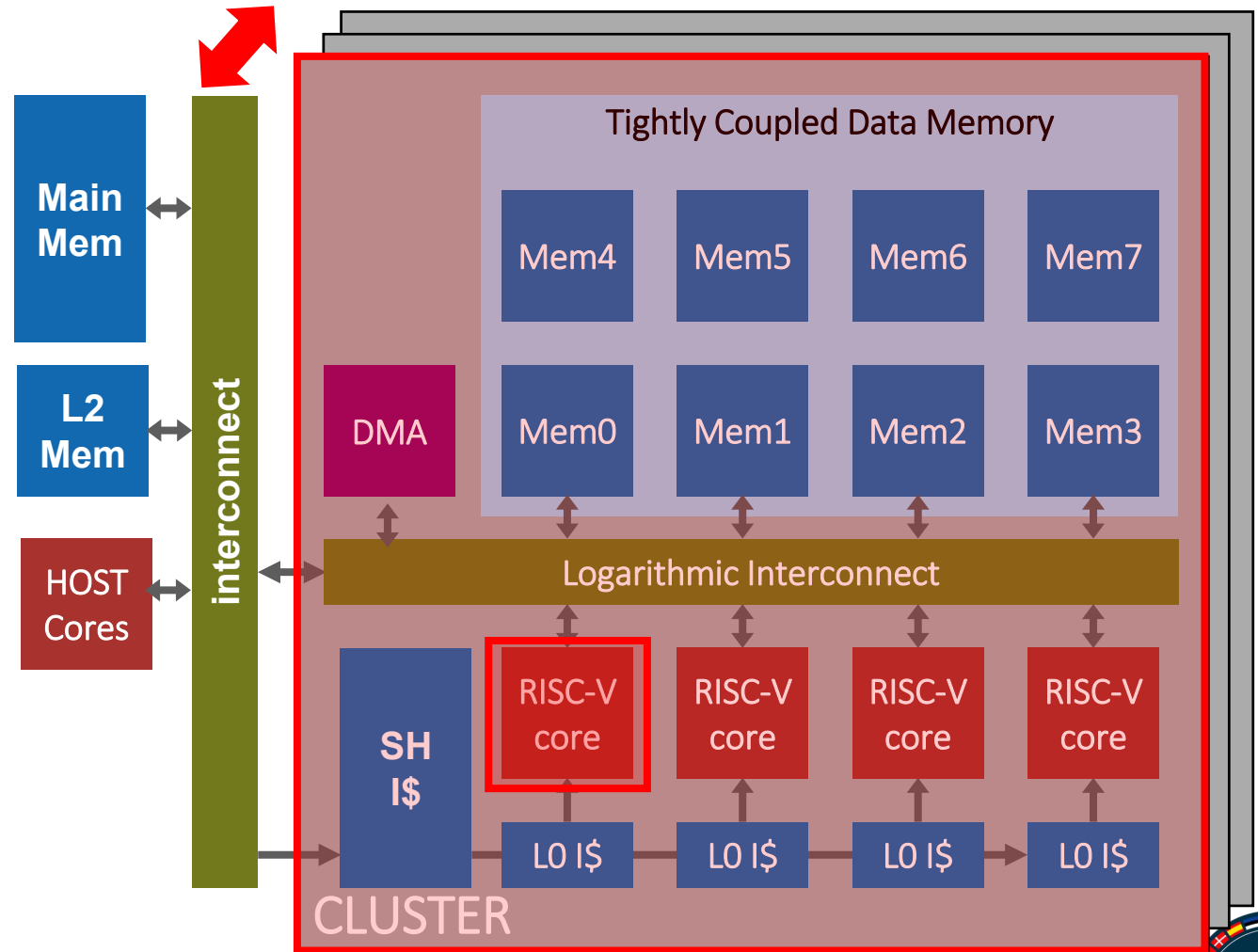
- Improving core efficiency with ISA and uAch extensions

- **Cluster**

- Efficient shared-mem cluster
- From a few to thousand processing element

- **Full platform**

- Heterogeneity: host, processor, accelerator
- One or multiple accelerators
- From chips to chiplets (2D to 3D)



PULP uses a **permissive** open source license

- All our development is on GitHub
 - HDL source code, testbenches, software development kit, virtual platform

<https://github.com/pulp-platform>



- Allows anyone to use, change, and make products without restrictions.

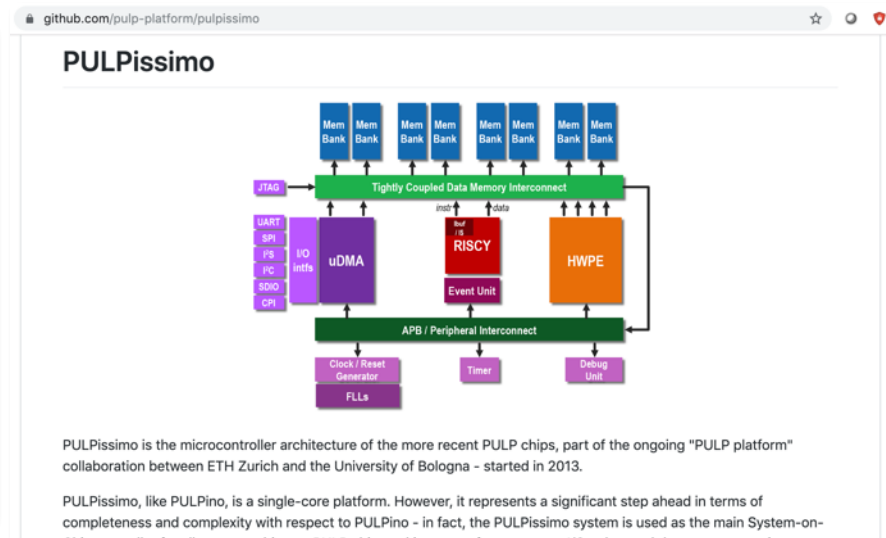
pulp-platform

Overview Repositories 217 Projects Packages People 12

Pinned

pulp Public
This is the top-level project for the PULP Platform. It instantiates a PULP open-source system with a PULP SoC (microcontroller) domain accelerated by a PULP cluster with 8 cores.
SystemVerilog ☆ 258 🍴 83

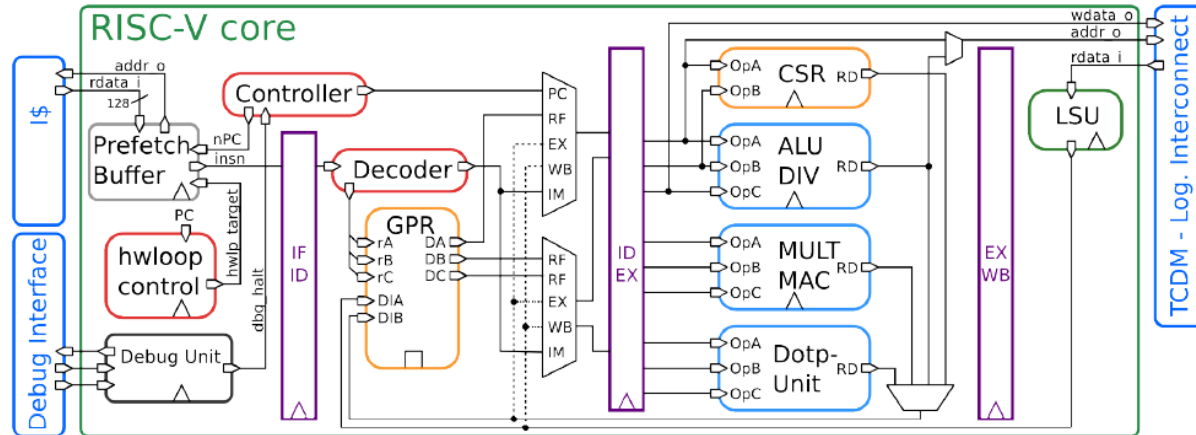
pulpissimo Public
This is the top-level project for the PULPissimo Platform. It instantiates a PULPissimo open-source system with a PULP SoC domain, but no cluster.
SystemVerilog ☆ 249 🍴 127



CV32E40P –aka RI5CY – RISC-V Core for edge AI



3-cycle ALU-OP, 4-cyle MEM-OP → IPC loss: LD-use, Branch



ISA is extensible *by construction* (great!)

```
for (i = 0; i < 100; i++)
    d[i] = a[i] + b[i];
```

Baseline 11 cyc./out

```
mv    x5, 0
mv    x4, 100
Lstart:
    lb    x2, 0(x10)
    lb    x3, 0(x11)
    addi  x10, x10, 1
    addi  x11, x11, 1
    add   x2, x3, x2
    sb    x2, 0(x12)
    addi  x4, x4, -1
    addi  x12, x12, 1
    bne   x4, x5, Lstart
```

Auto-incr 8 cyc./out

```
mv    x5, 0
mv    x4, 100
Lstart:
    lb    x2, 0(x10!)
    lb    x3, 0(x11!)
    addi  x4, x4, -1
    add   x2, x3, x2
    sb    x2, 0(x12!)
    bne   x4, x5, Lstart
```

HW Loop 5 cyc./out

```
lp.setupi 100, Lend
    lb    x2, 0(x10!)
    lb    x3, 0(x11!)
    add   x2, x3, x2
Lend:    sb x2, 0(x12!)
```

SIMD 1.25 cyc./out

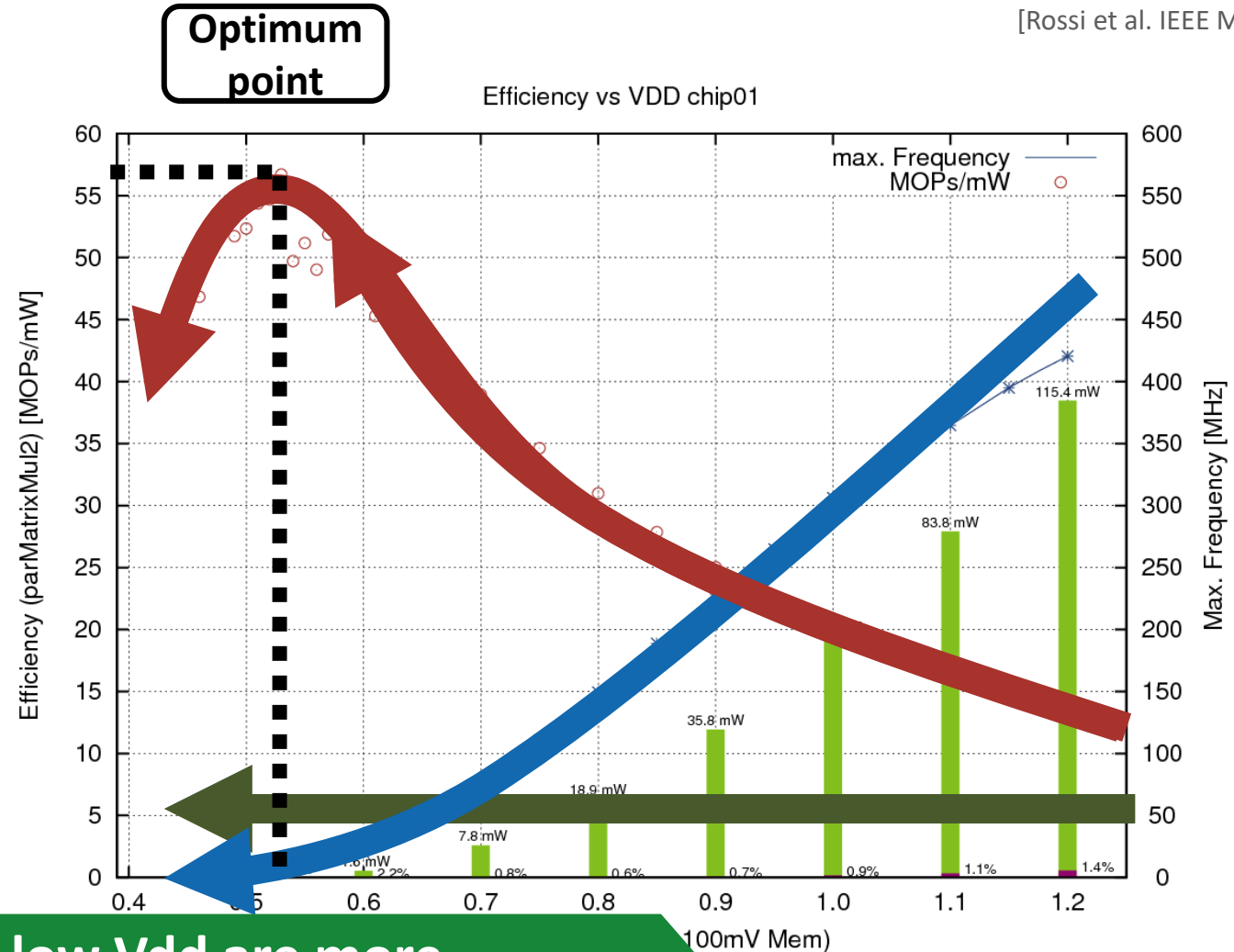
```
lp.setupi 25, Lend
    lw    x2, 0(x10!)
    lw    x3, 0(x11!)
    pv.add.b x2, x3, x2
Lend:    sw x2, 0(x12!)
```

Scaling performance: the case for parallelization



[Rossi et al. IEEE Micro 2017]

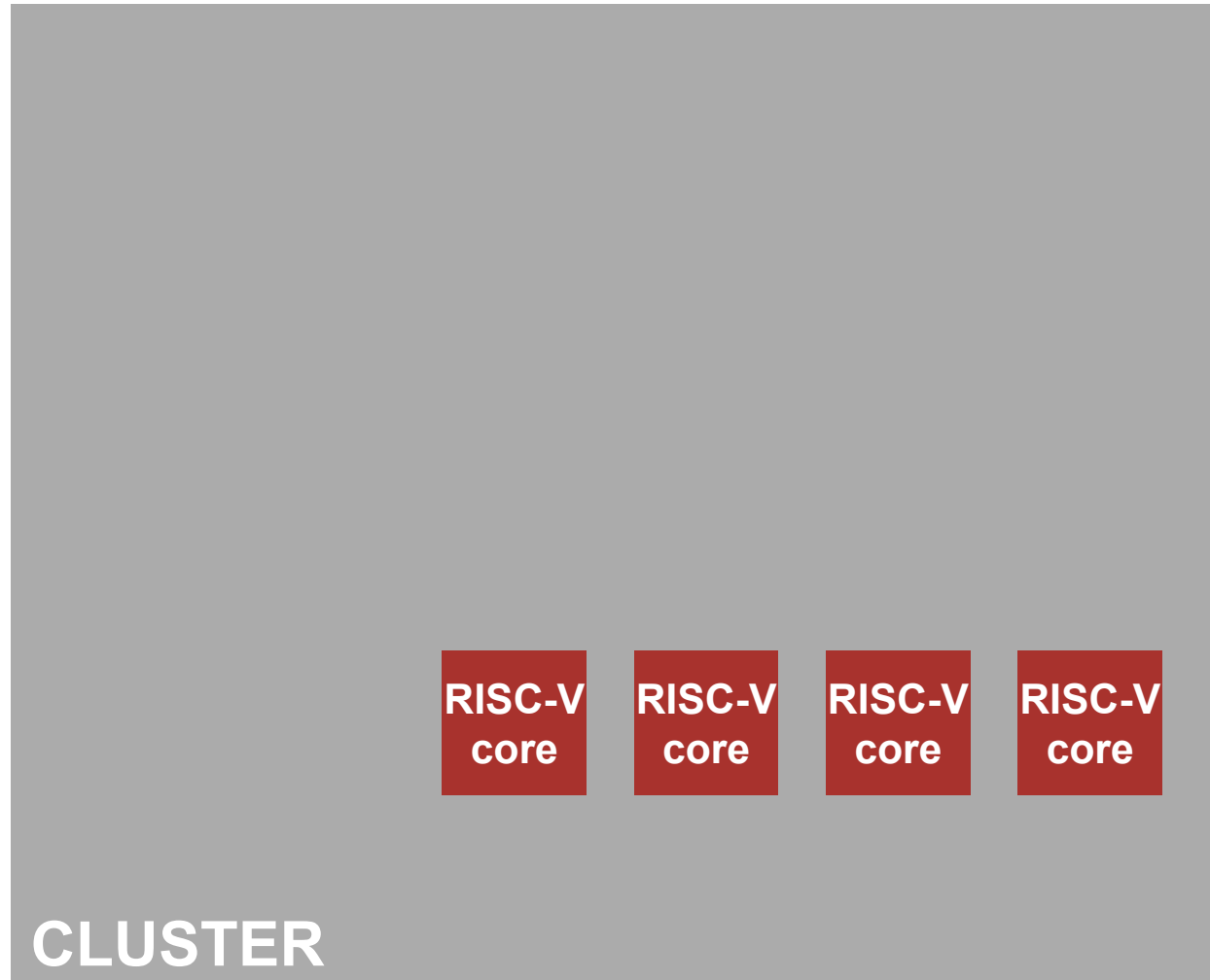
- As **VDD** decreases, **operating speed** decreases as well.
- However **efficiency** increases → more work done per Joule
 - Until leakage effects start to dominate
- Put more units in parallel to get performance up and keep them busy with a parallel workload



N cores running at moderate f, low Vdd are more energy efficient than a single core at $N \times f$, high Vdd



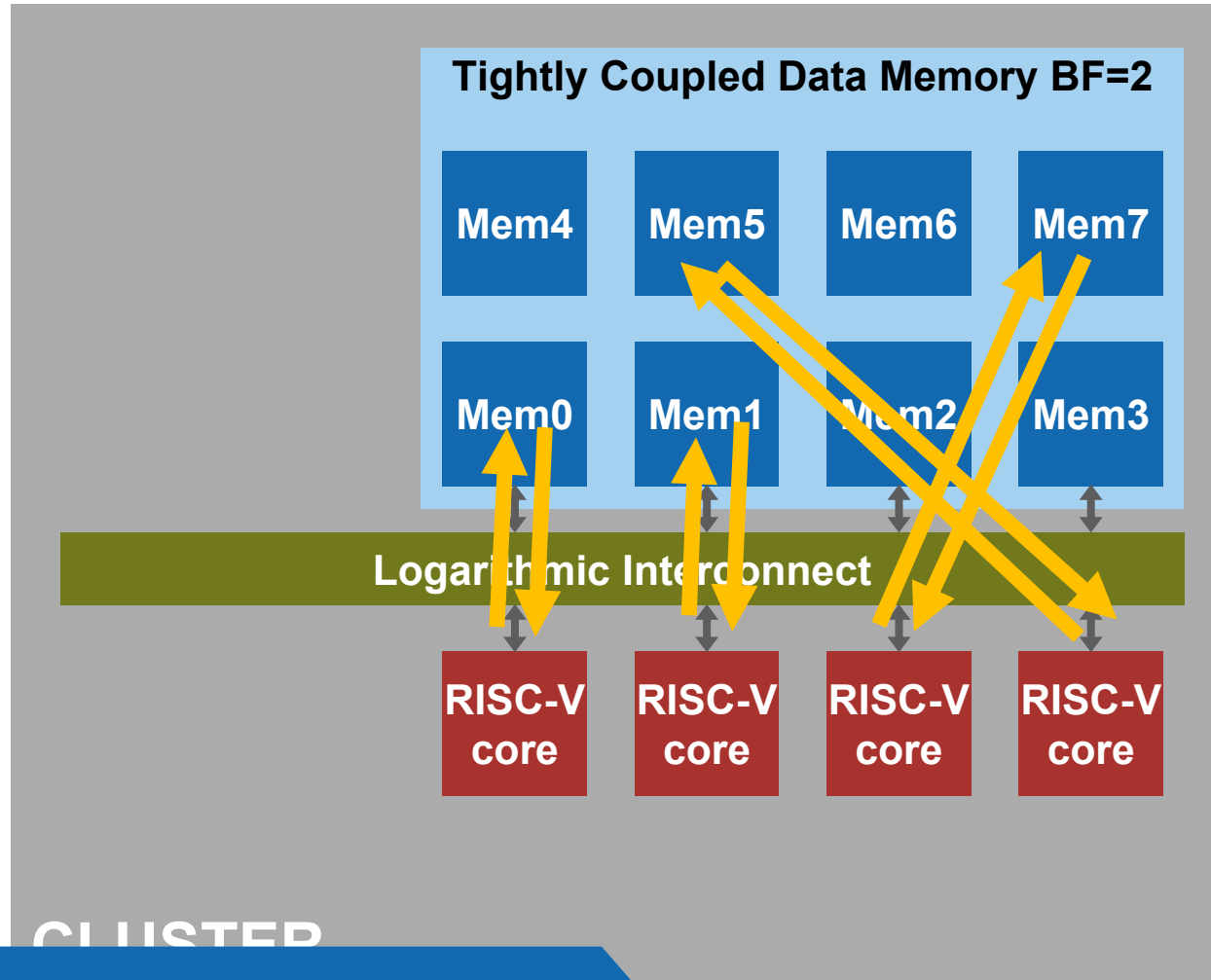
Let's design a cluster with multiple (4-16) RISC-V cores



Low-Latency shared Tightly Coupled Data Memory (TCDM)



- Parallel memory access with low contention
 - Multi-banked, address-interleaved L1
- Fast interconnect with physical design awareness
 - Logarithmic depth of combinational switchboxes

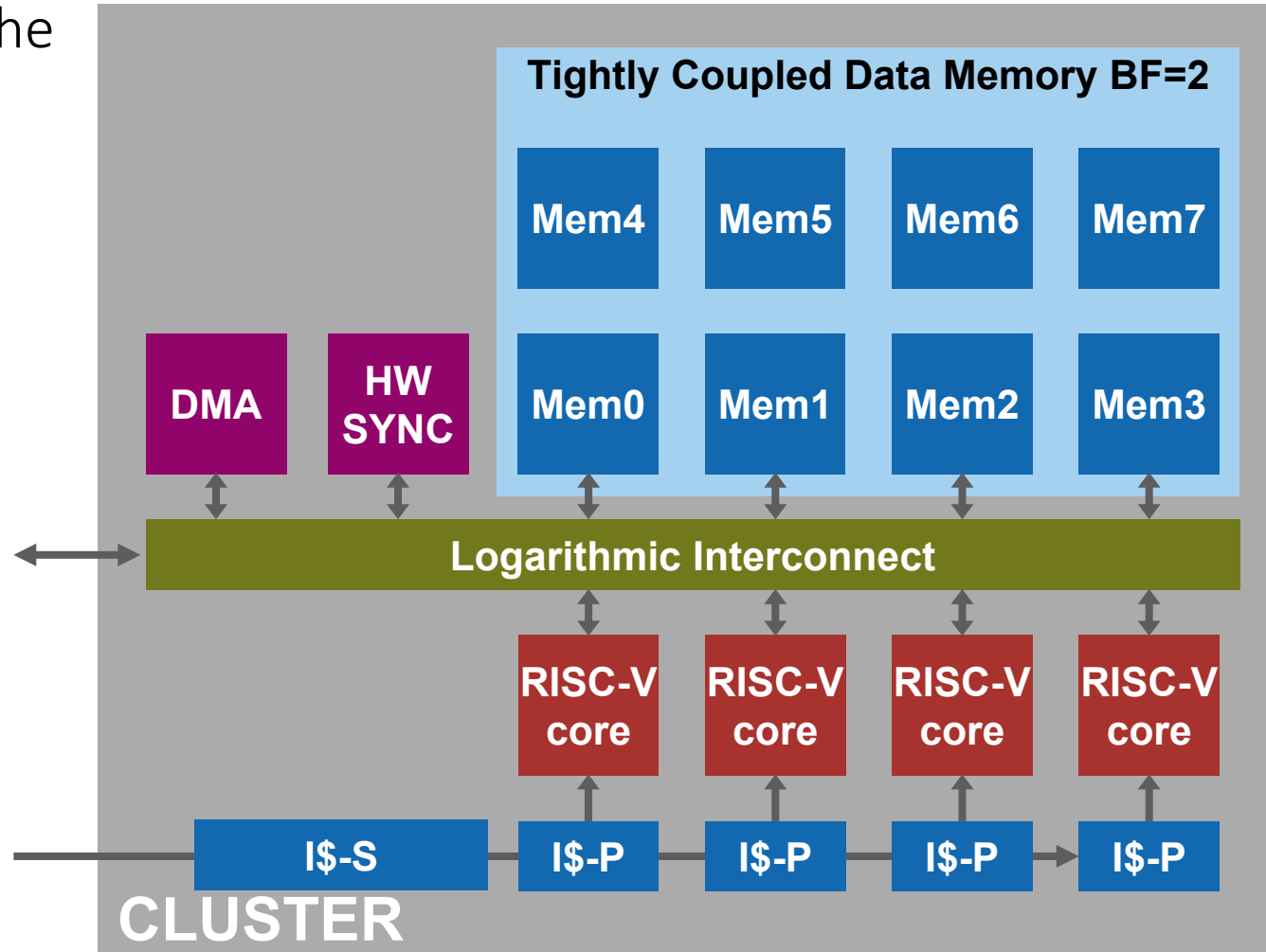


Trade-off between memory size and latency



Full Cluster

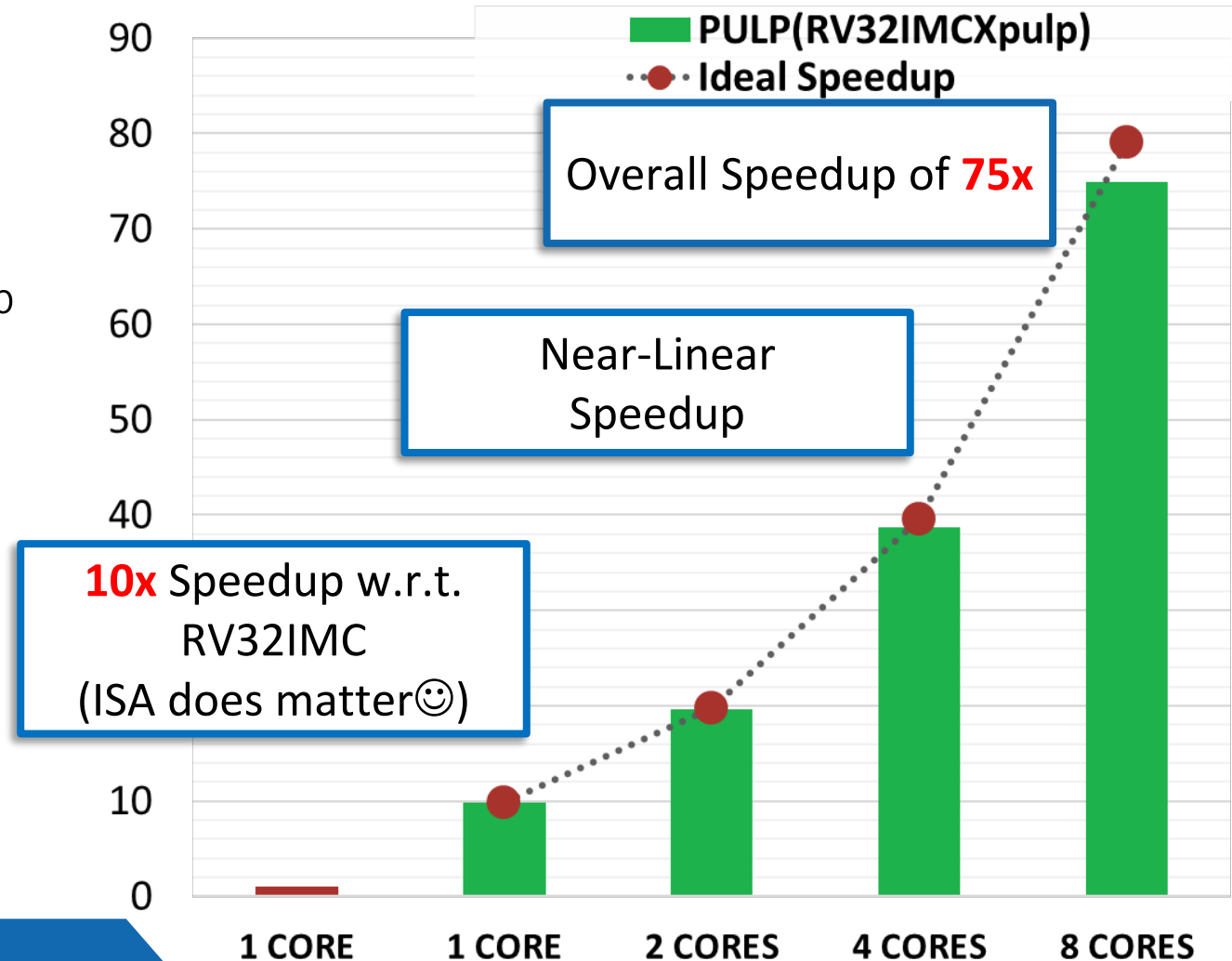
- Hierarchical Instruction Cache
 - Speeds up instruction fetching
- DMA
 - Facilitates memory transfers
- Event Unit
 - Facilitates core synchronization
- Ports to a host



Combining ISA extension + Efficient parallel execution



- 8-bit convolution
 - Open source DNN library
- 10x through xPULP
 - Extensions bring real speedup
- Near-linear speedup
 - Scales well for regular workloads
- 75x overall gain
- 7-8 GMACs
 - 250MHz
 - 4 MAC/Cycle (8bit)



More GOPS for less power

[Garofalo et al. Philos. Trans. R. Soc 20]



Overview

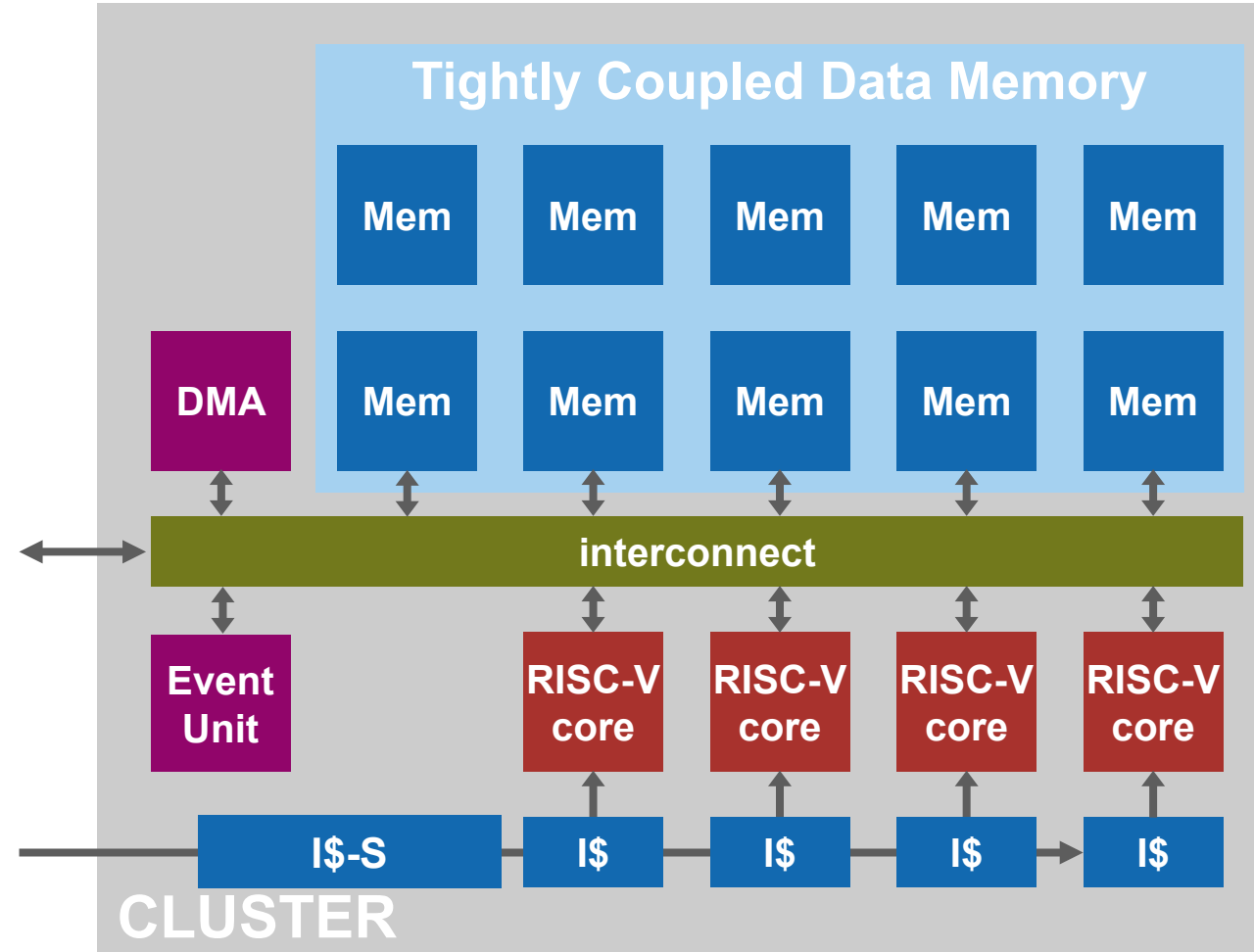


1. The PULP Platform
2. Leveraging the PULP cluster for a Fault-Tolerant System
3. Using Heterogeneity for High-Performance Systems

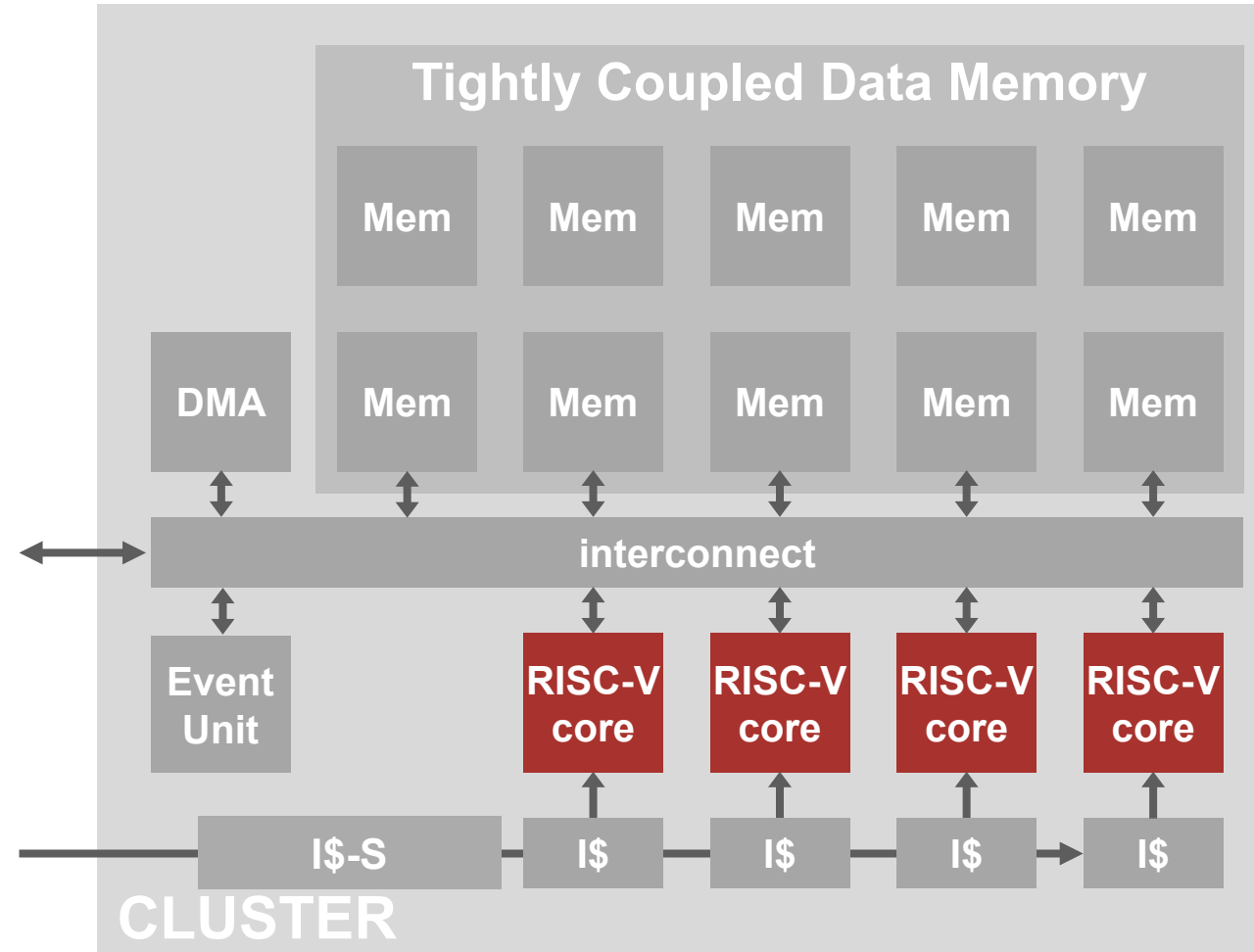
Towards a space-proof PULP cluster

- Cores
- Memory
- Interconnect, DMA, EU

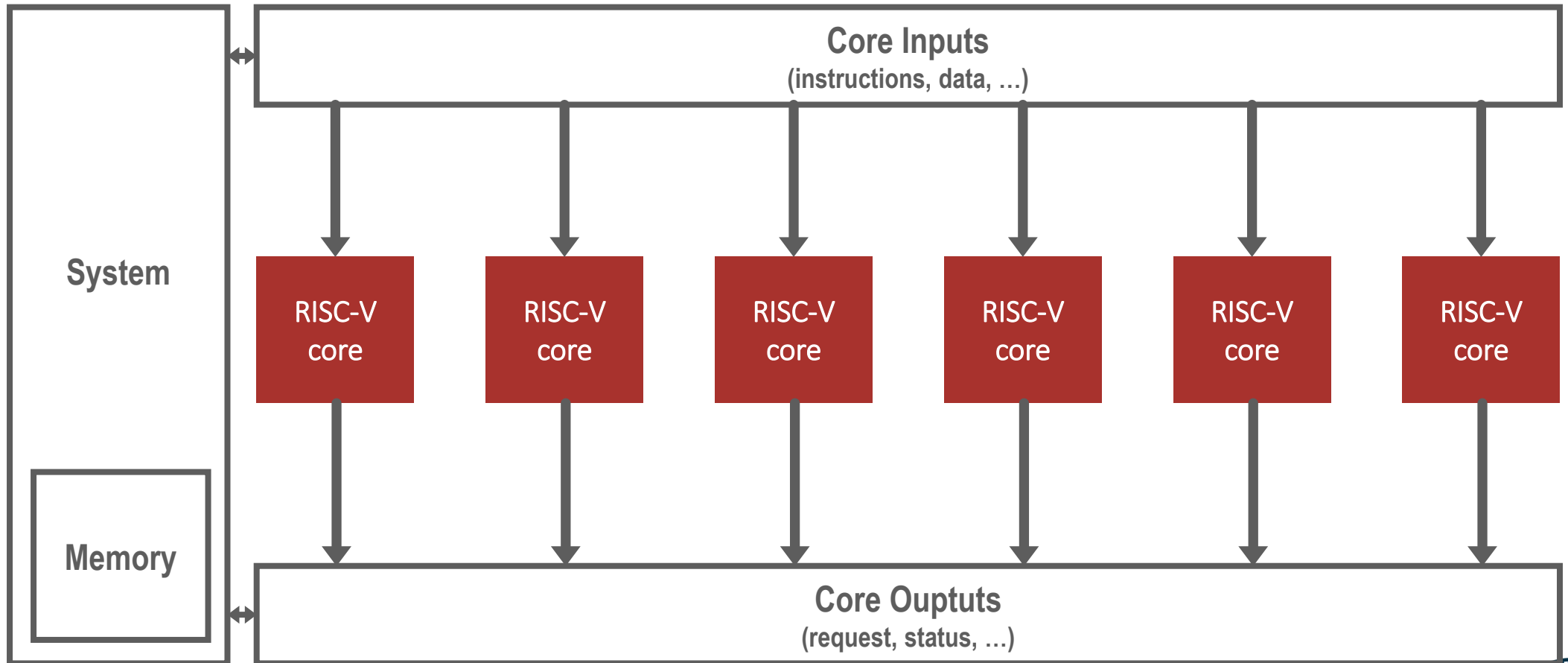
Component-based
vs. Architectural approach



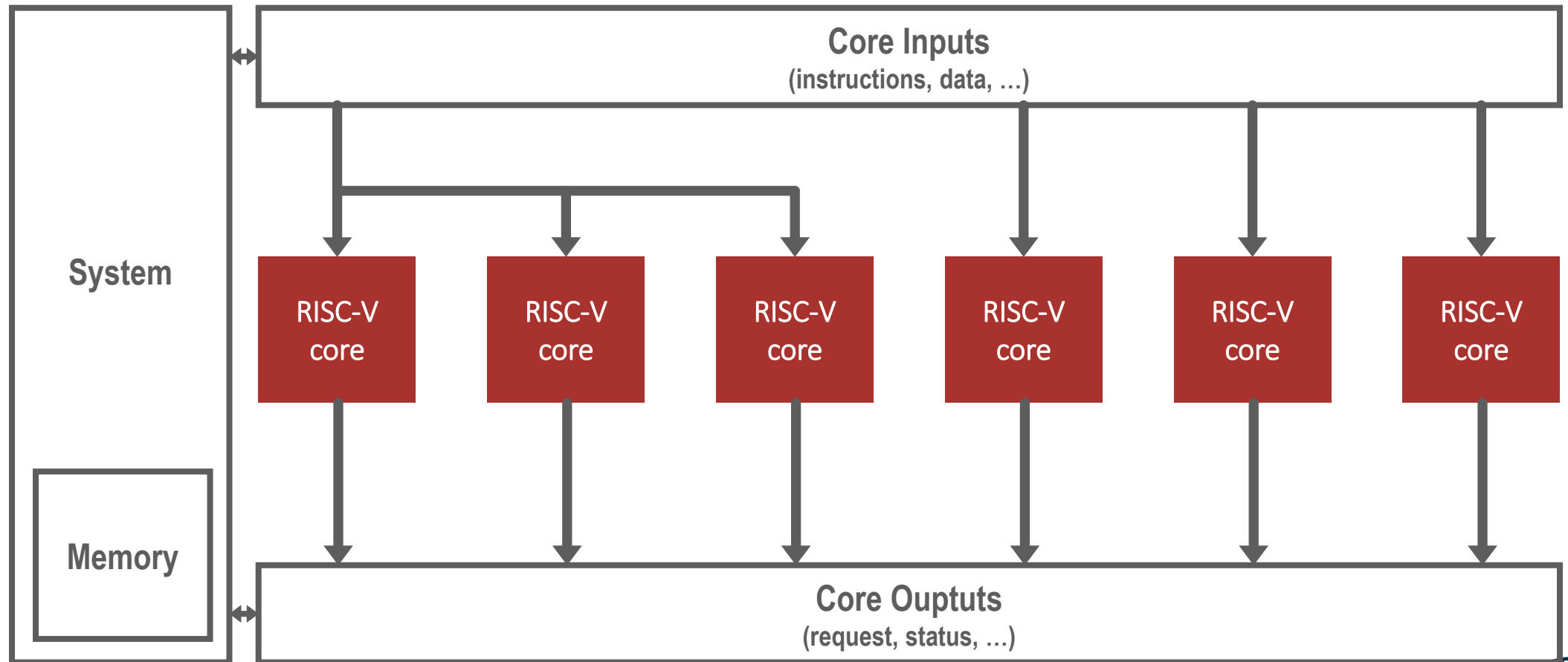
Leveraging Parallelism



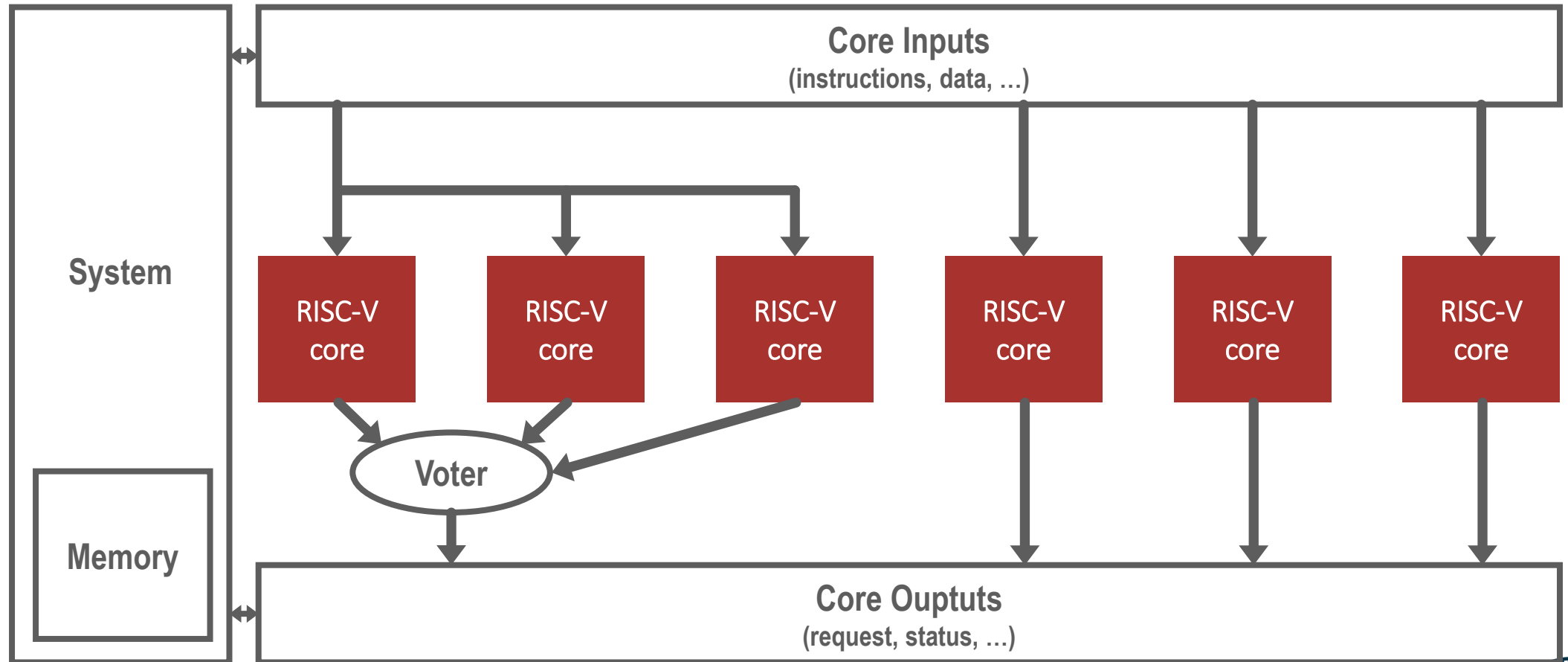
Protecting the Cores



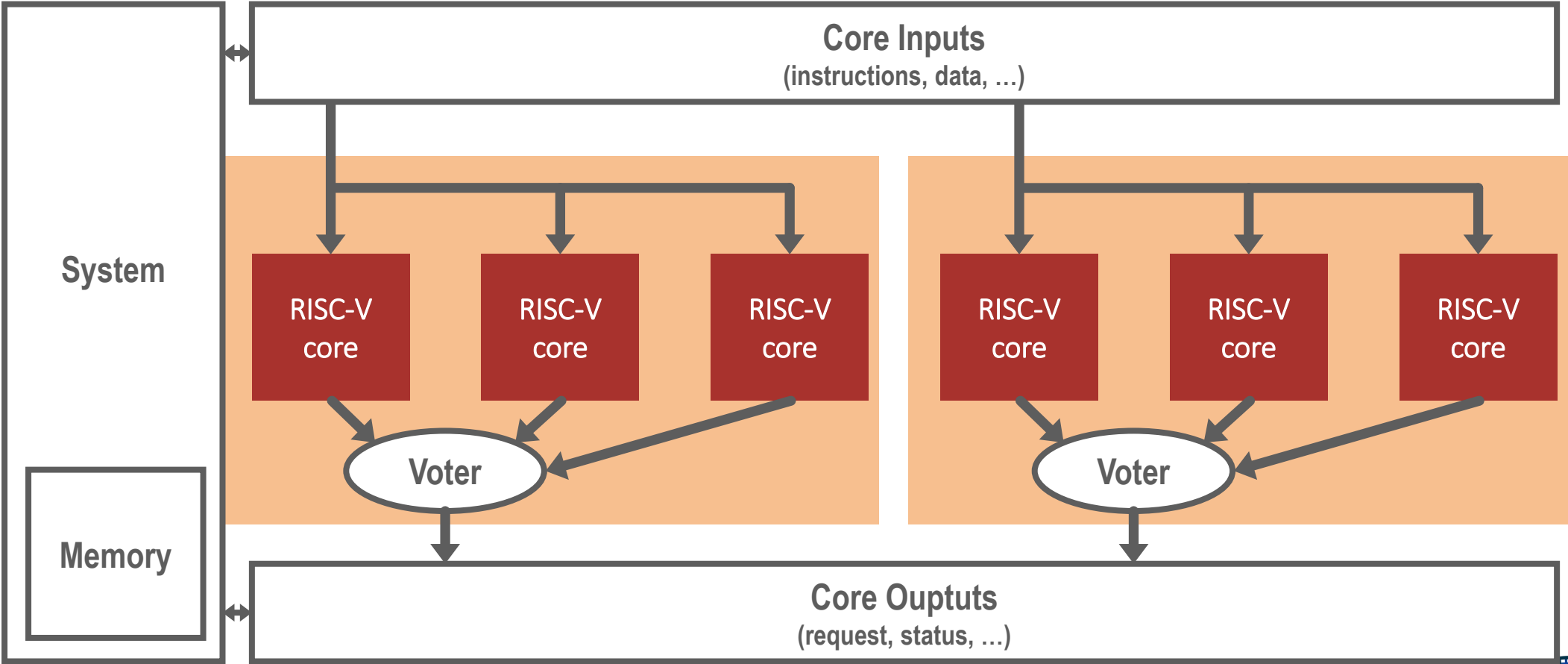
Protecting the Cores



Protecting the Cores

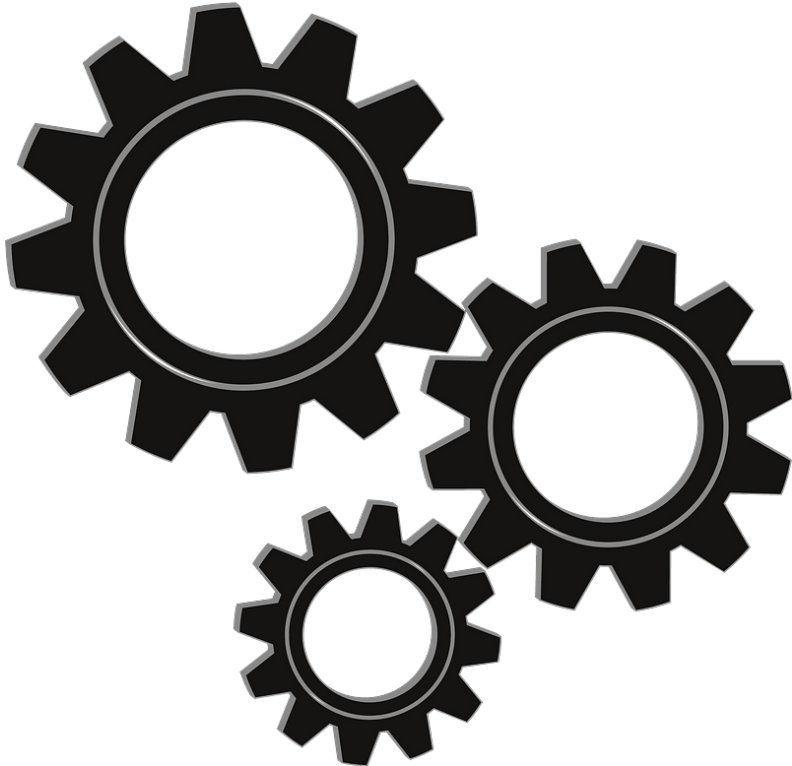


Triple-Core Lock Step

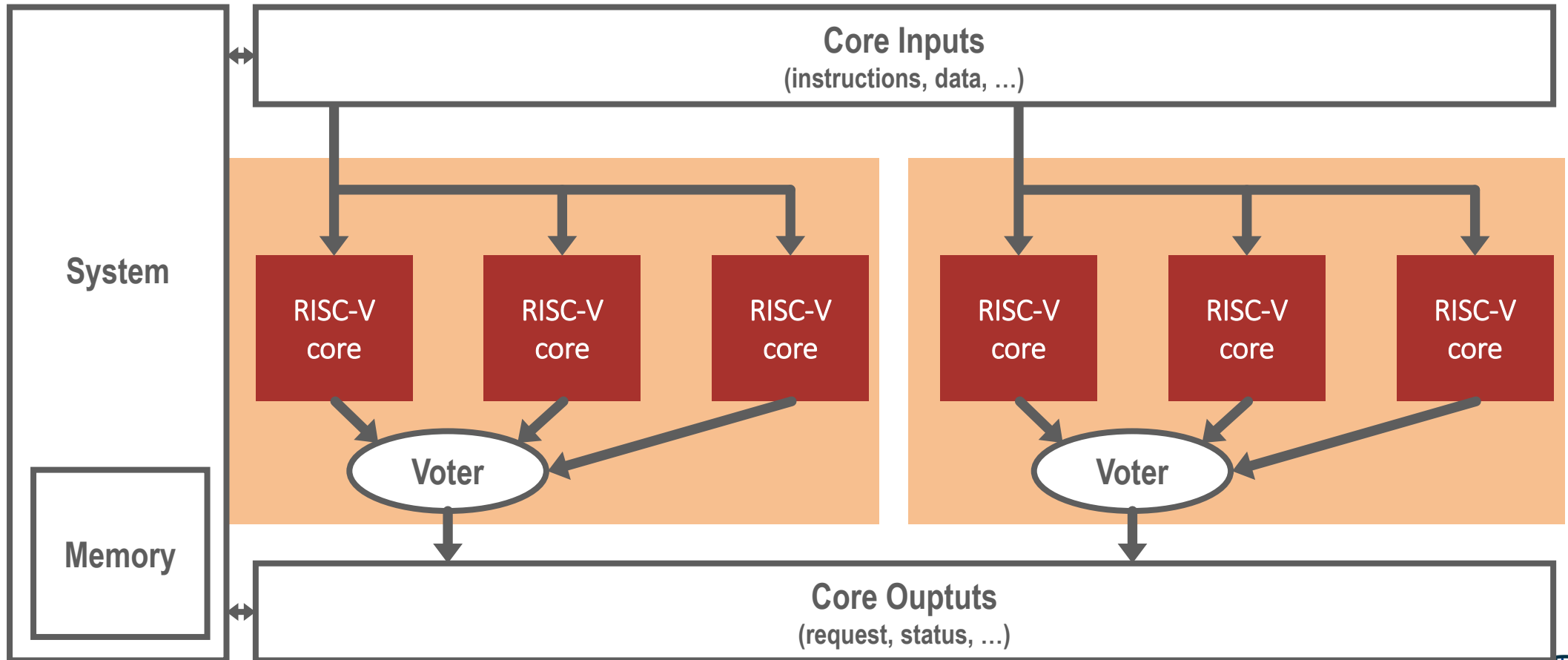


Redundancy Requirement

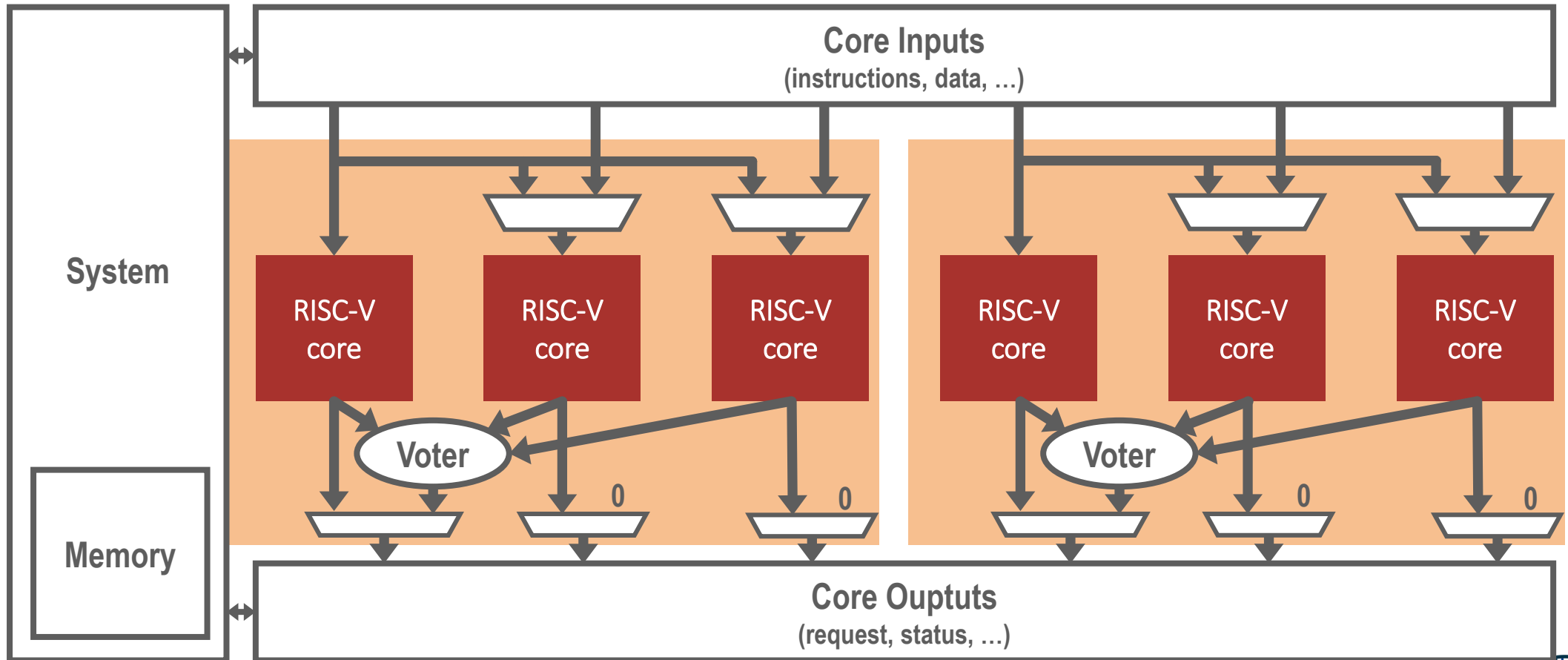
- Mission-Critical Application
 - Must be correct
- Data Processing Algorithm
 - Individual errors can be tolerated



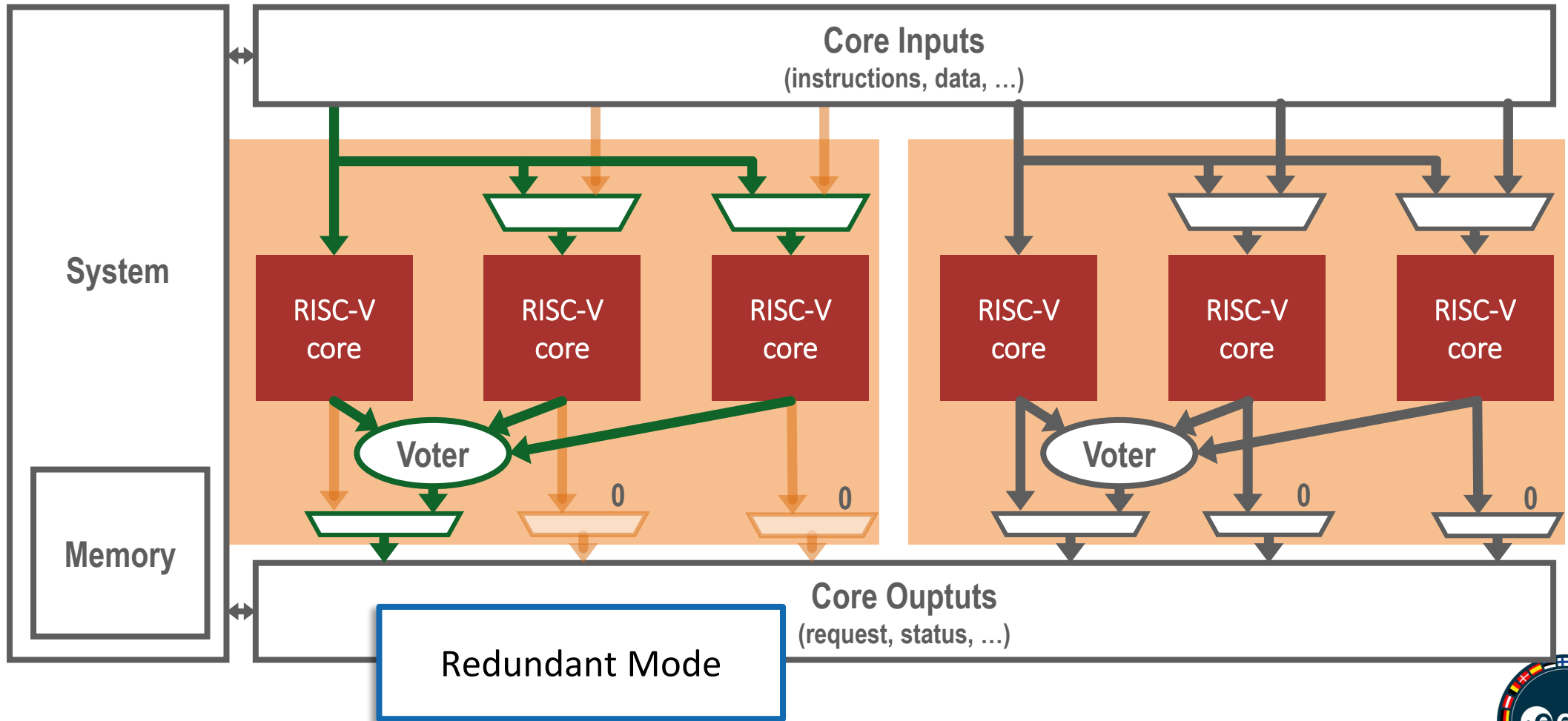
On-Demand Redundancy Grouping



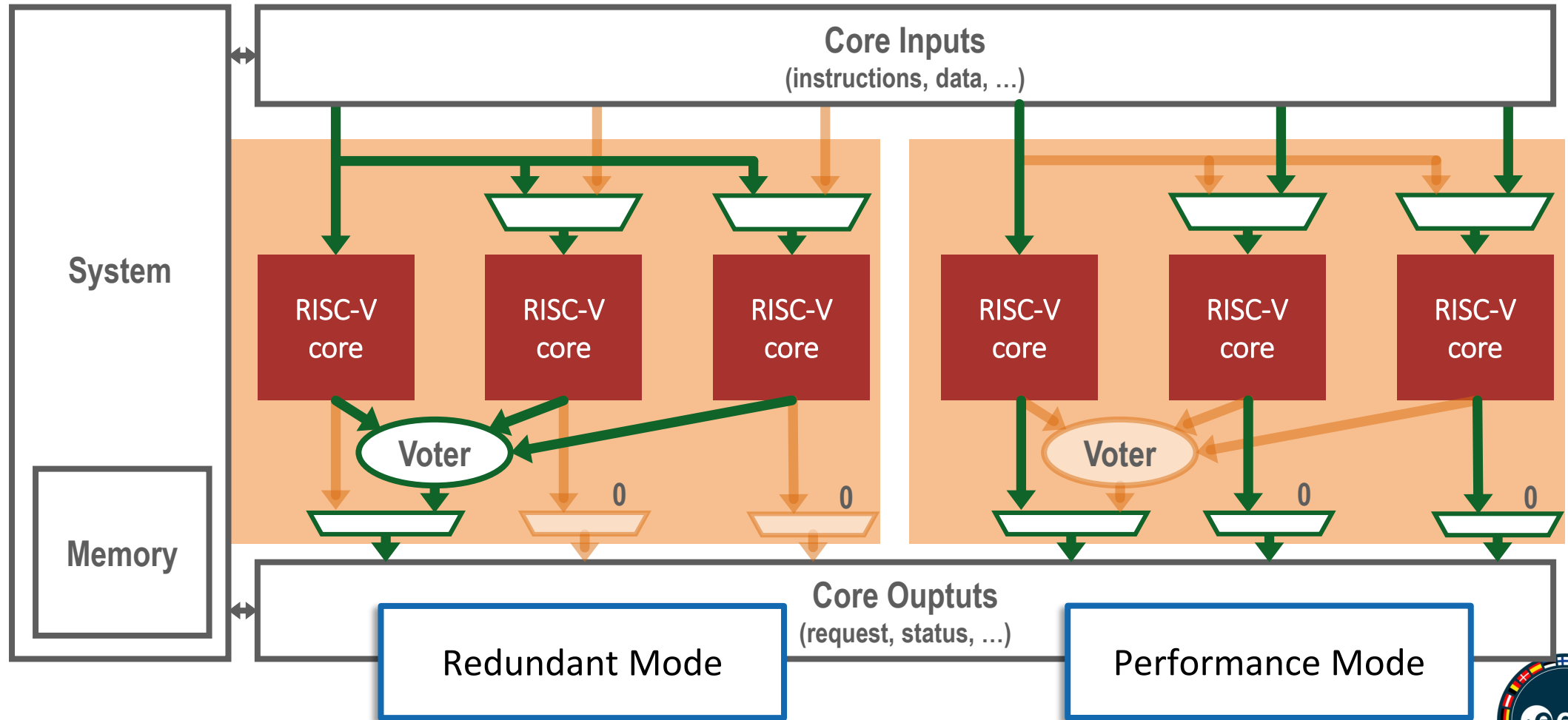
On-Demand Redundancy Grouping



On-Demand Redundancy Grouping



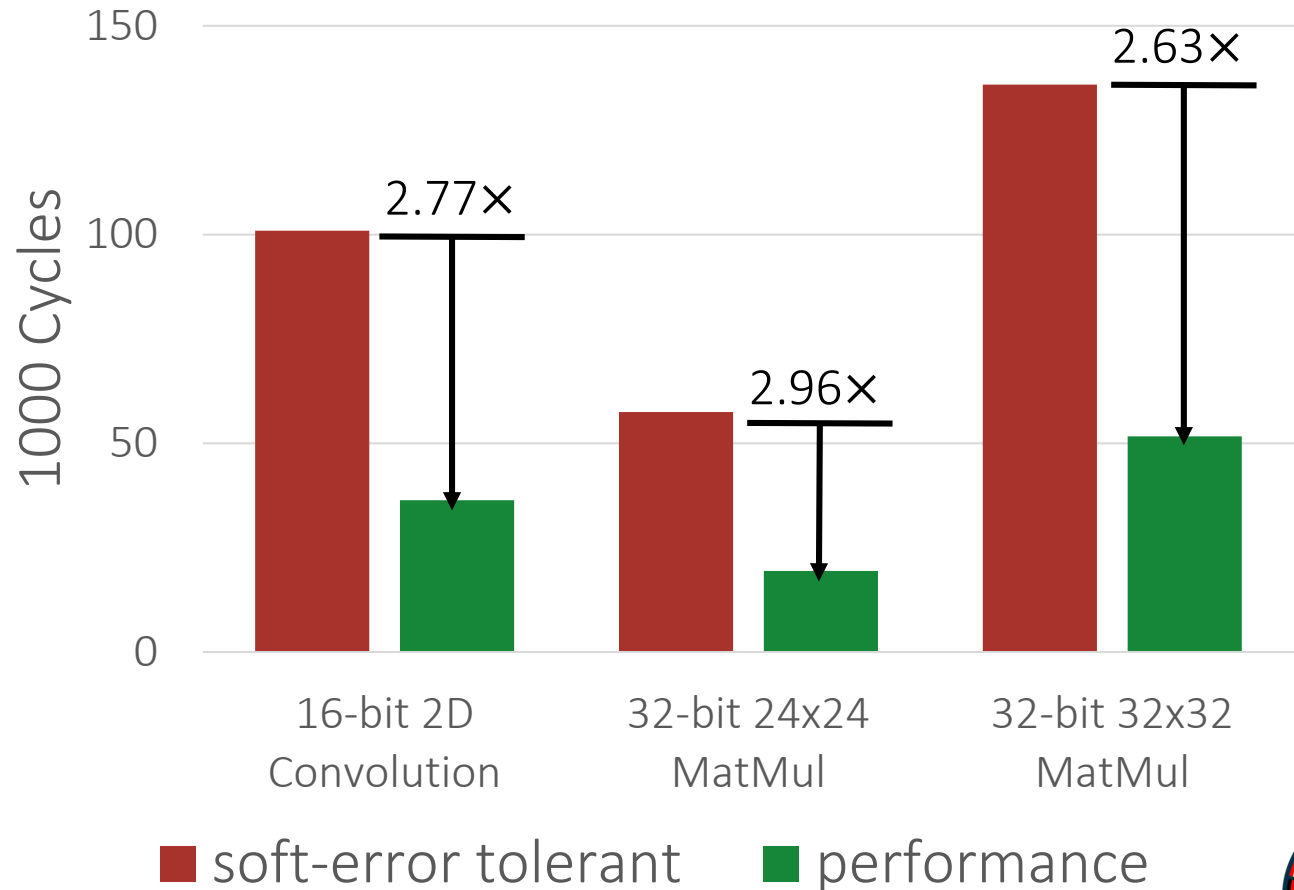
On-Demand Redundancy Grouping



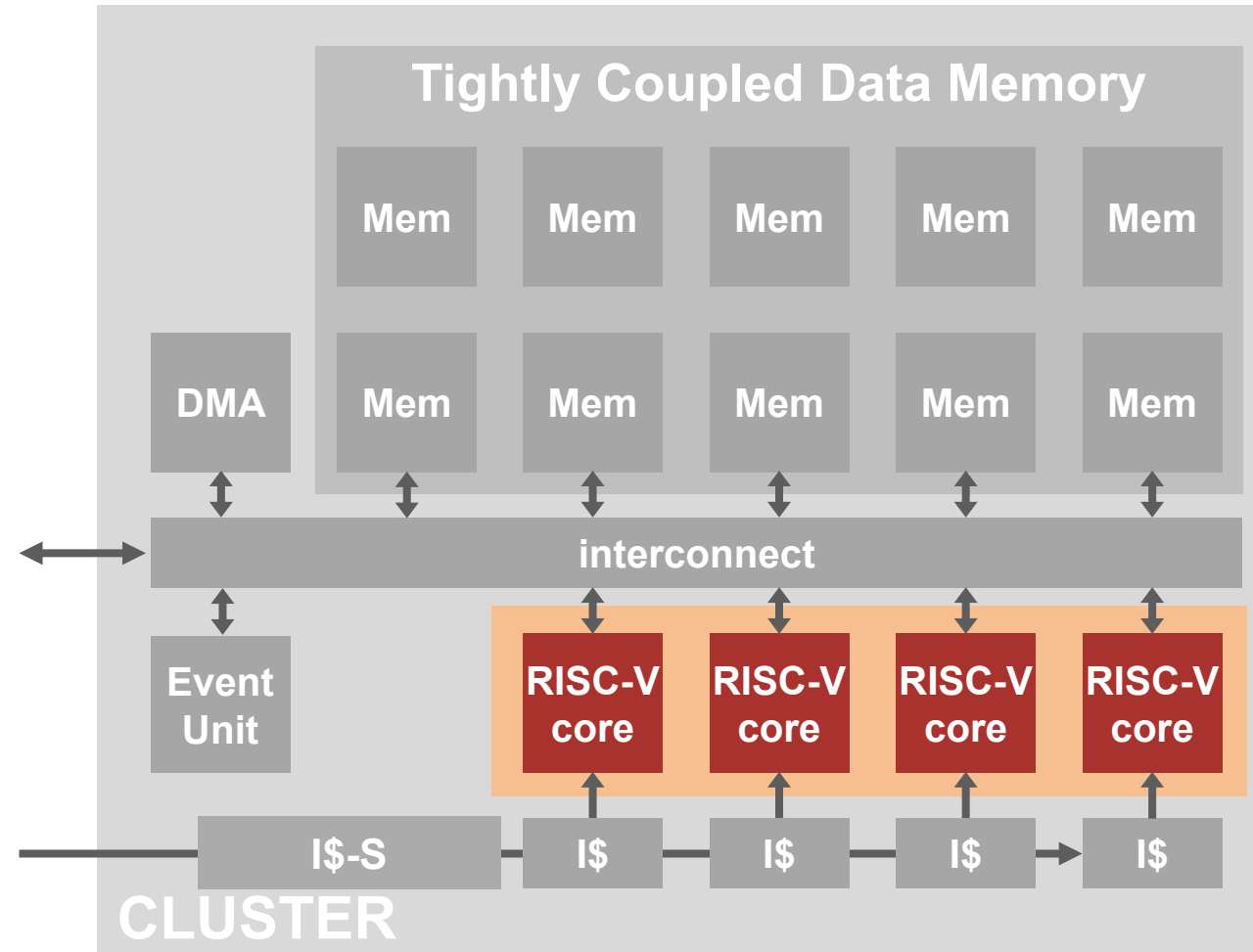
On-Demand Redundancy Grouping



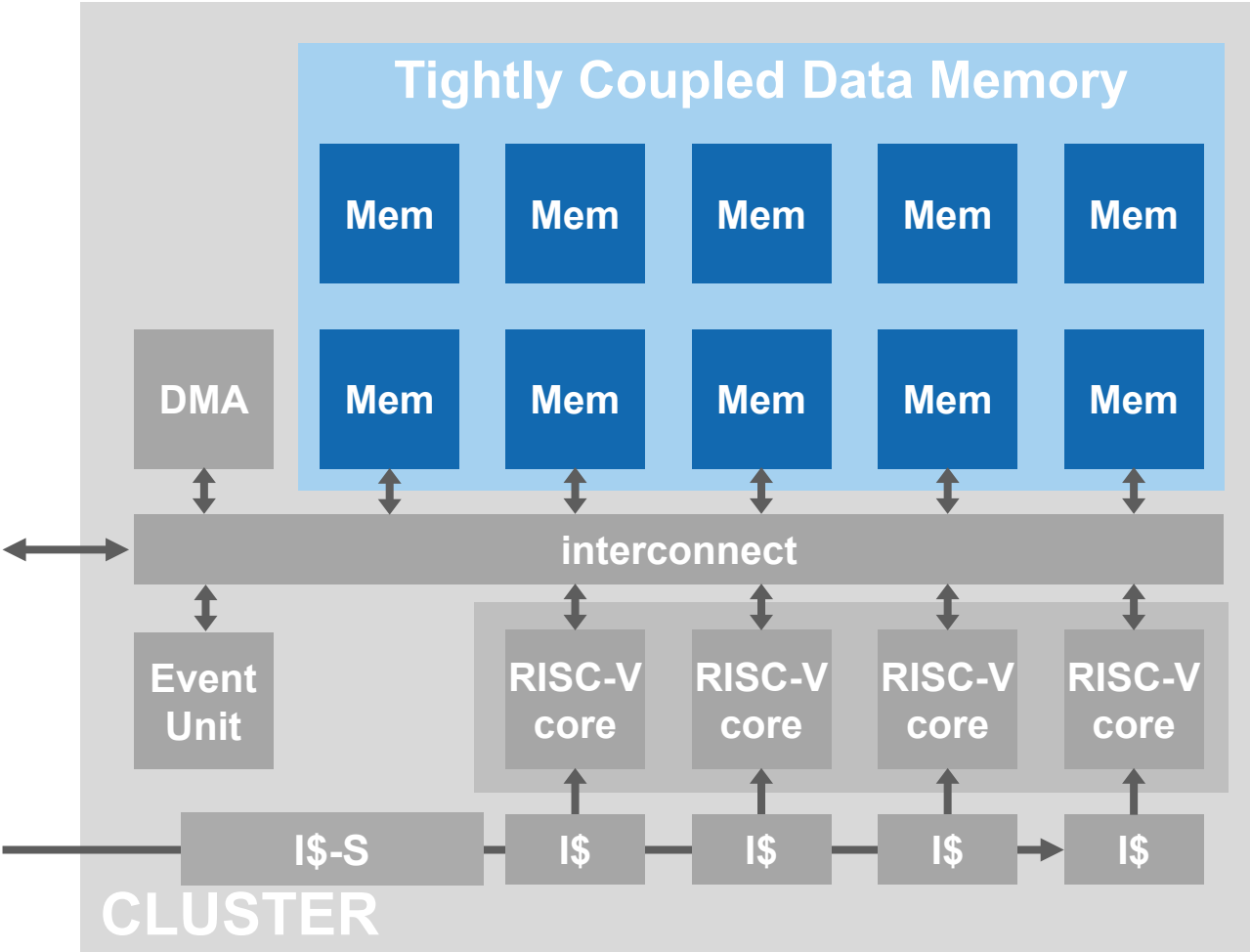
- Up to **3x speedup** for non-critical tasks
- **<60'000 cycles** to switch modes
- **<1% area overhead** for the cluster
 - ODRG for 3 cores requires ~11% area of a single core



From Cores...

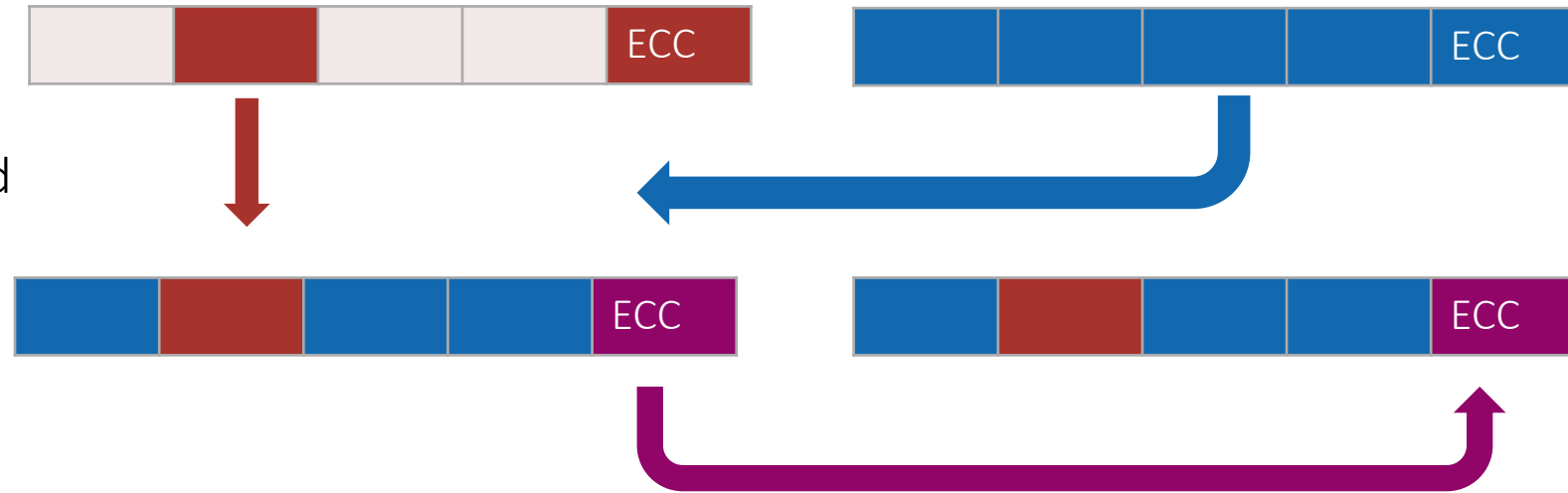


To Memory

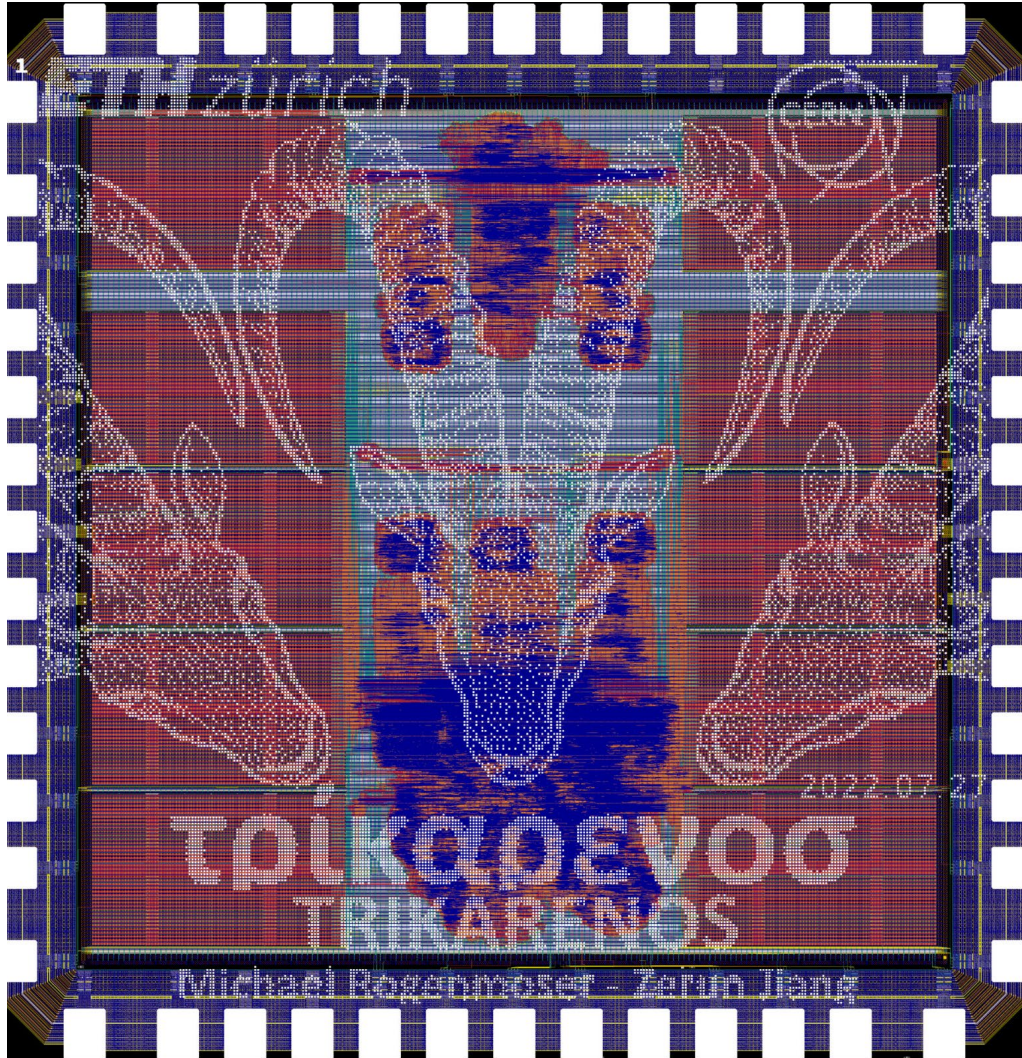


ECC Memory

- 39-32 bit Hsiao Code
- Efficient byte-store
 - Coremark has <10% non-word stores
 - <1% overhead
- Scrubber per memory bank



Trikarenos



- Fault tolerance features implemented in an ASIC
- TSMC 28nm demonstrator
- Incorporates testing features
- Planned to fly on ARIS' CubeSat

aris
space to grow



Overview



1. The PULP Platform
2. Leveraging the PULP cluster for a Fault-Tolerant System
3. Using Heterogeneity for High-Performance Systems

Heterogeneous + Parallel... Why?



- Processors can do two kinds of useful work:

Decide (jump to different program part)

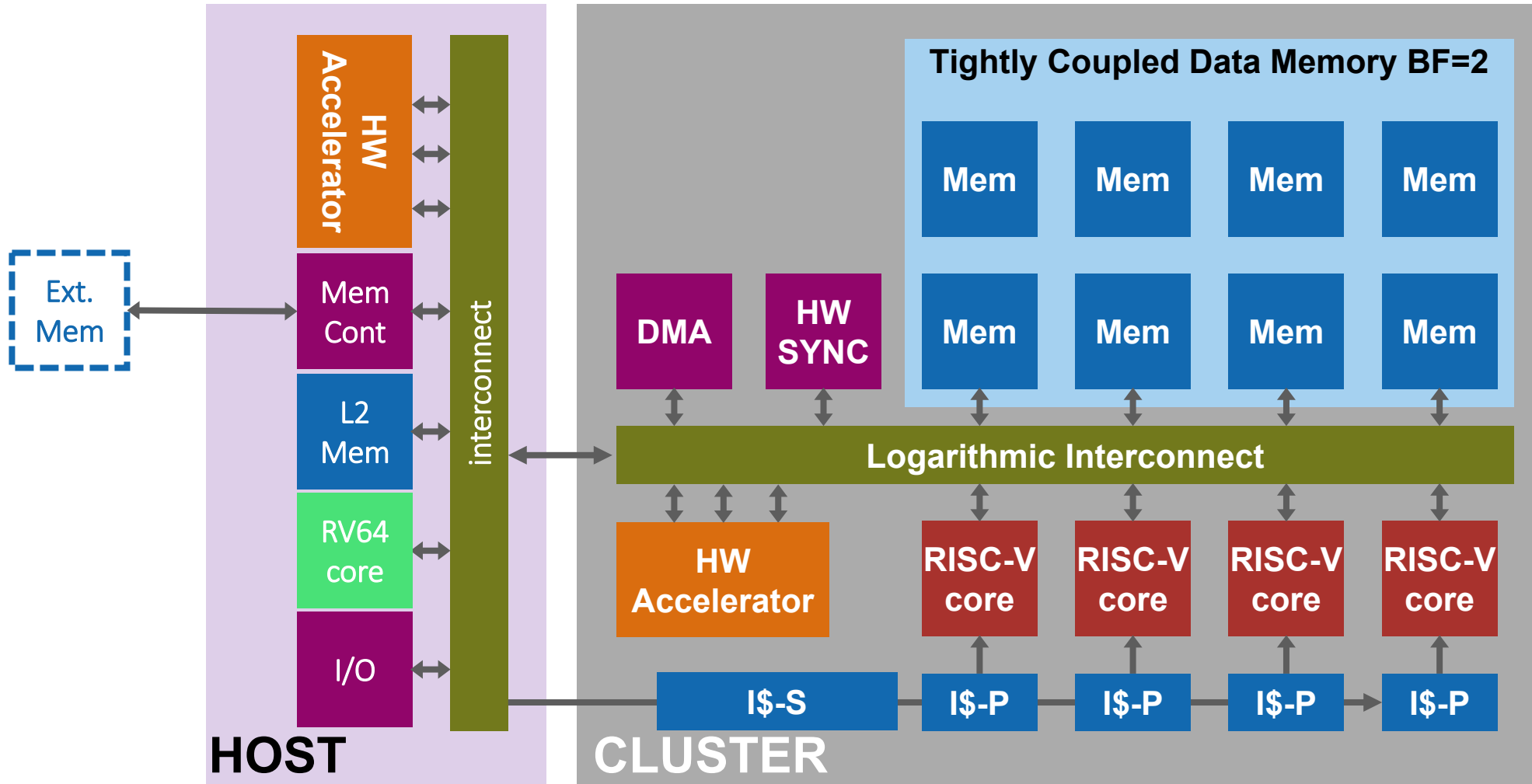
- Modulate flow of **instructions**
- Mostly sequential decisions:
 - Don't work too much
 - Be clever about the battles you pick (latency is king)
- **Lots of decisions**
Little number crunching

Compute (plough through numbers)

- Modulate flow of **data**
- Embarassingly data parallel:
 - Don't think too much
 - Plough through the data (throughput is king)
- **Few decisions**
Lots of number crunching

- Today's workloads are dominated by "Compute":
 - Tons of data, few (as fast as possible) decisions based on the computed values,
 - Data-Oblivious Algorithms (ML, or better DNNs are so!)
 - Large data footprint + sparsity

From Cluster to full Platform

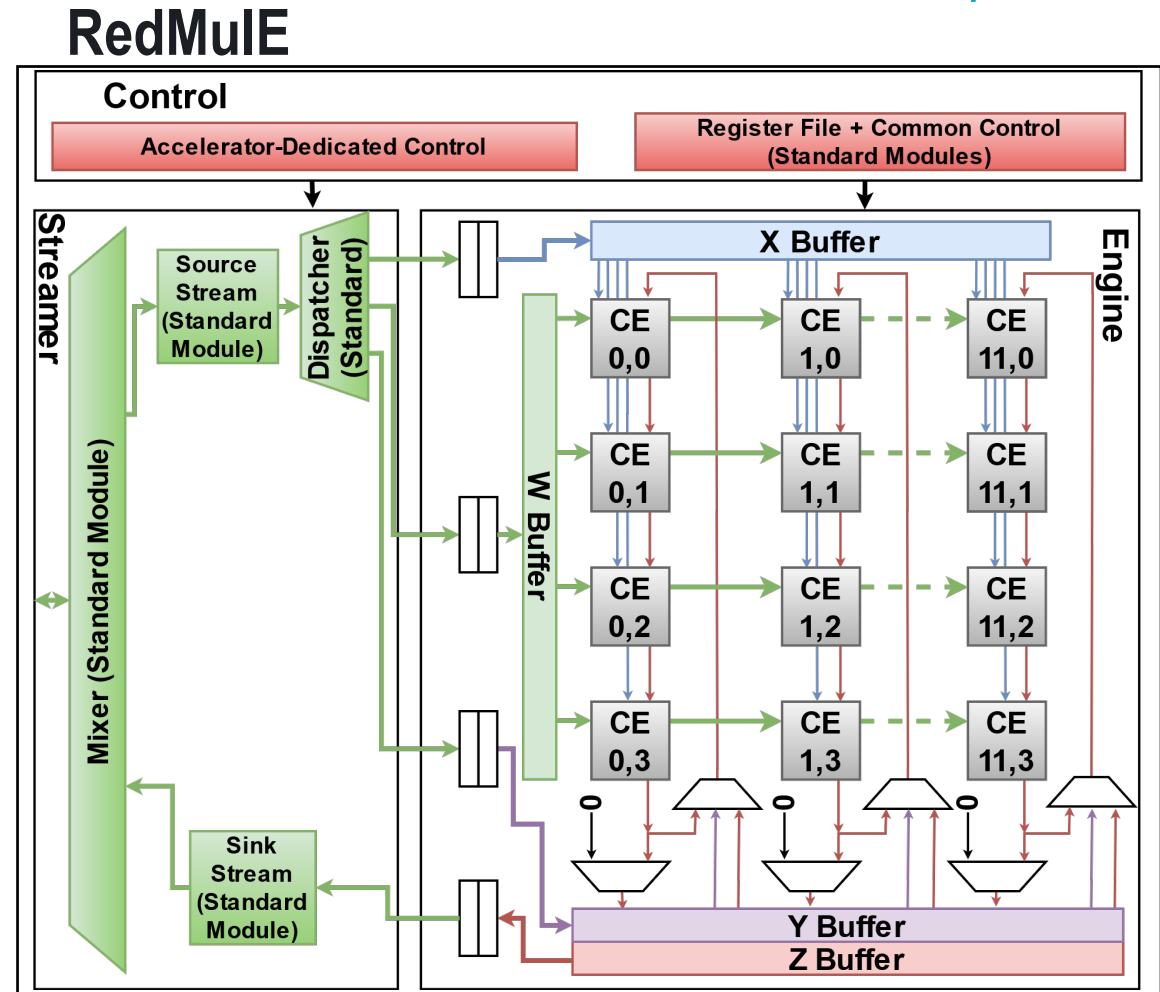


HW Accelerator Case Study: RedMulE



Reduced-Precision Matrix Multiplication Engine

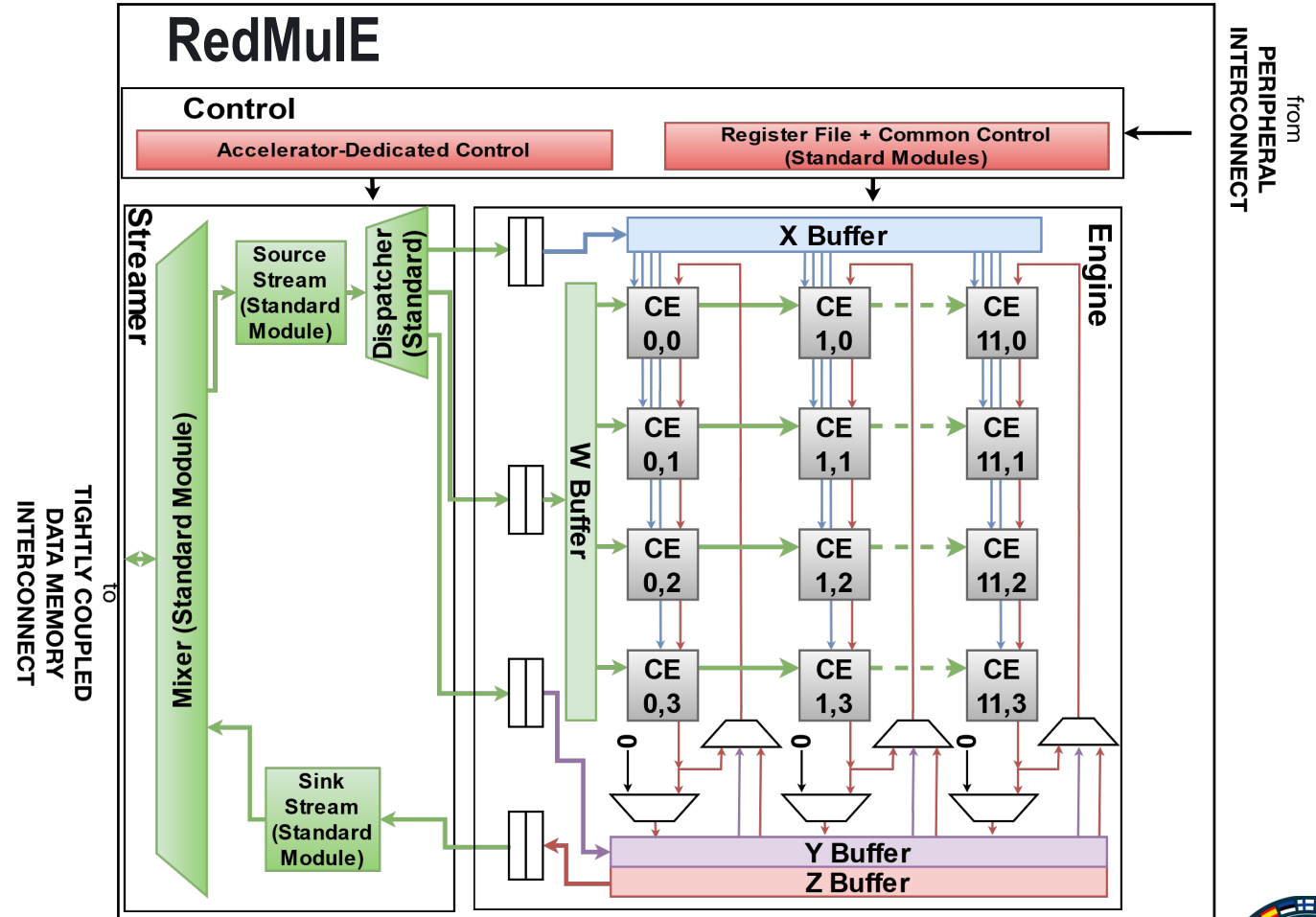
- Accelerates general matrix-matrix operations (**GEMM-Ops**) using mixed-precision floating point (**FP16/FP8**) arithmetic
- Enables **on-chip training** of **generalized** deep learning models
- Accelerates **graph** and **control-flow-based** algorithms (all-pairs shortest paths, minimum spanning tree, and others)



HW Accelerator Case Study: RedMule



- Flexible template for hardware accelerators design
- Interface designed by composition of open-source IPs for **control** + **memory** interfaces through *HWPE-streams* → specialized streamers convert from streams to mem (TCDM)
- Memory access capabilities equivalent or superior to SW (misaligned access, delayed stores, high bandwidth) on part or all of the memory map



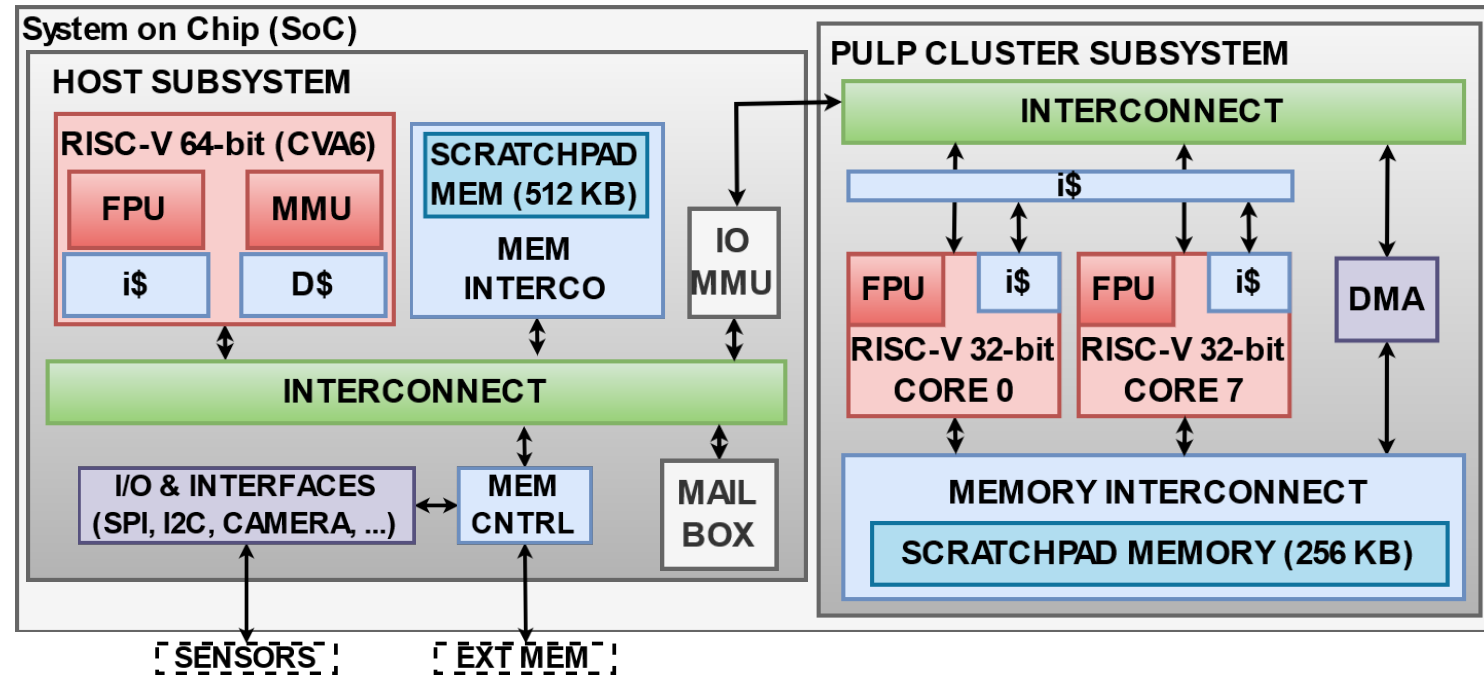
Heterogeneous SoC



- 64-bit **Linux-capable** RISC-V core
- Shared memory among the two existing domains
- Interaction with **external** environment (memory, IO)
- **Cross-domain** communication: mailbox + TLB for virtualization

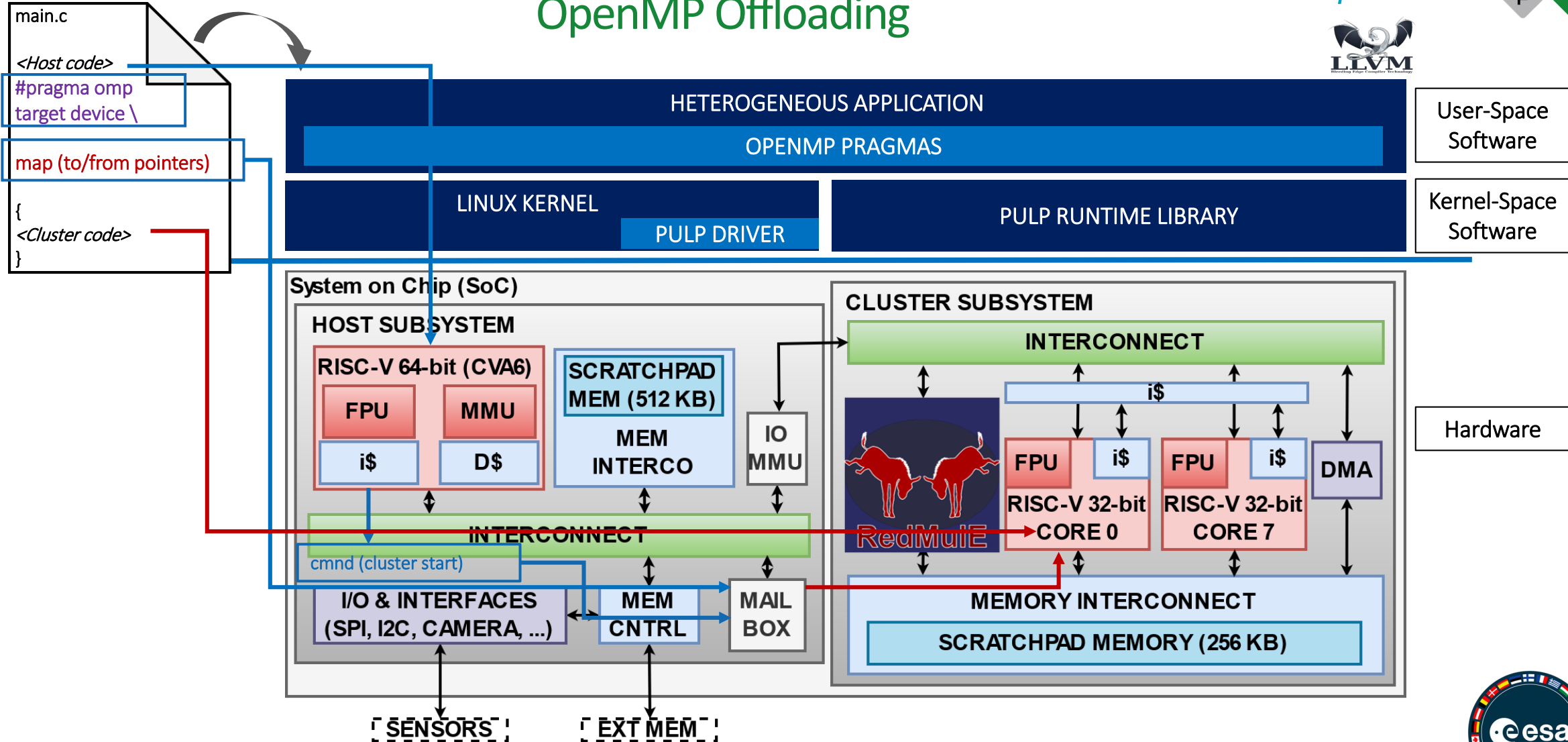
Features

- **Ease of programmability** of Single-Board Computers (**Linux**)
- **Performance** and **efficiency** of parallel multi-core cluster



Heterogeneous SoC for Space Applications

OpenMP Offloading



Heterogeneous SoC for Space Applications

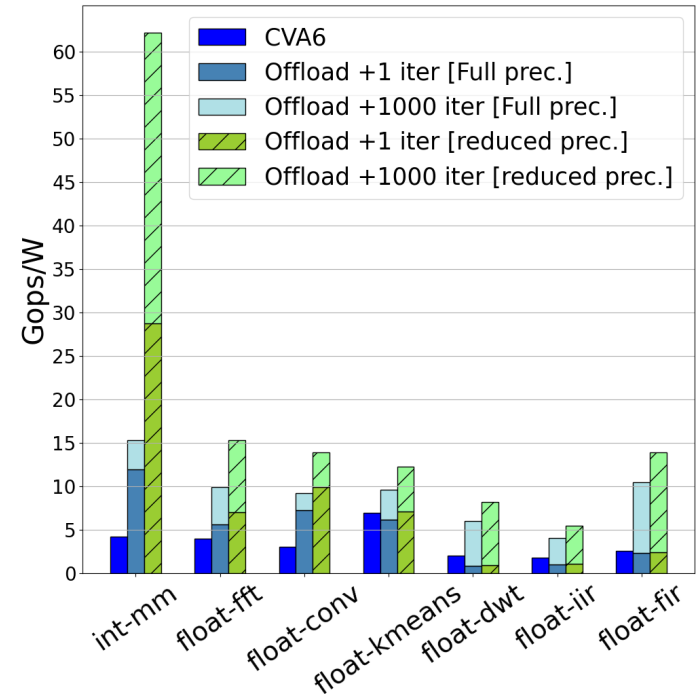
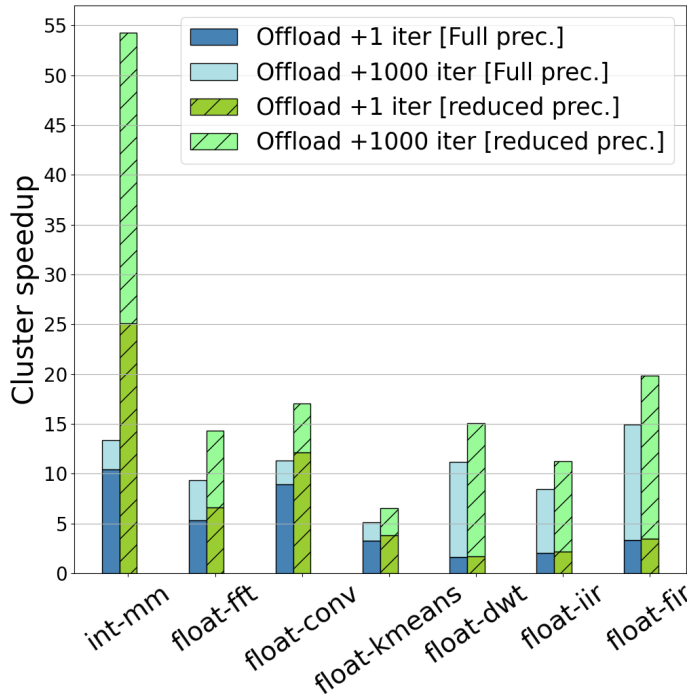
Hybrid DSP Kernel-Set (Integer + FP)

Performance (@ 500 MHz):

- Integer: Cluster speedup up to **55x** with respect to Host
- FP: Cluster speedup up to **20x** with respect to Host

Energy Efficiency (22 nm, 0,8 V)

- Integer: up to **62 GOPS/W** (12,7x better than Host)
- FP: up to **15 GFLOPS/W** (3,19x better than Host)



Future Developments



- Hybrid Modular Redundancy (**HMR**) in PULP clusters: **selectable** dual or triple core **lockstep**
- **Host** Core microarchitecture **protection** (redundancy, error correction)
- Reliable **on-chip interconnect**
- Reliable **peripherals communication** (interfaces redundancy)
- Reliability of **DMA** and **Hardware Accelerators**

Conclusion: What is PULP?



- **Core**

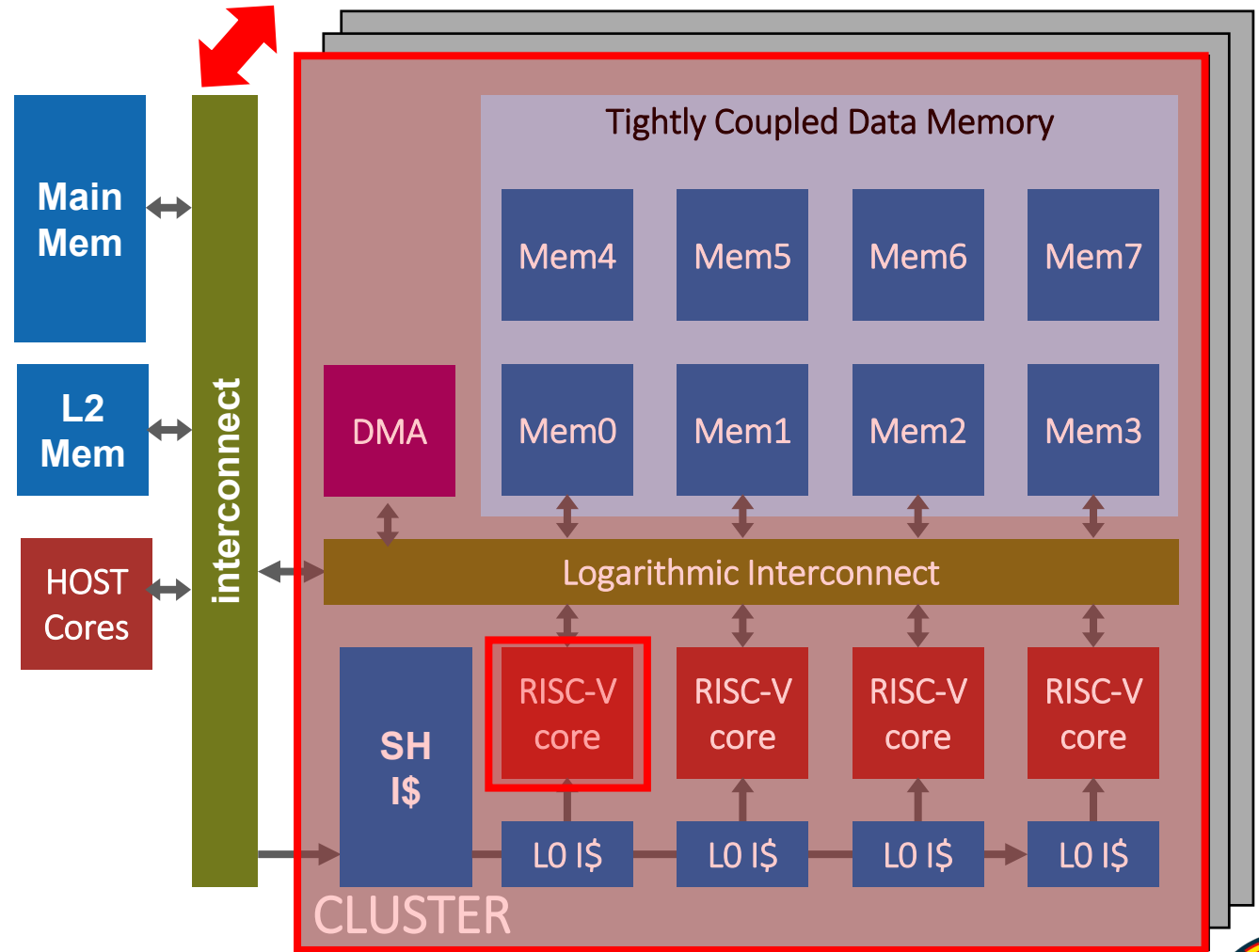
- Improving core efficiency with ISA and uAch extensions

- **Cluster**

- Efficient shared-mem cluster
- From a few to thousand processing element

- **Full platform**

- Heterogeneity: host, processor, accelerator
- One or multiple accelerators
- From chips to chiplets (2D to 3D)





PULP

Parallel Ultra Low Power

Luca Benini, Ahmad Mirsalari, Alessandro Capotondi, Alessandro Nadalini, Alessandro Ottaviano, Alessio Burrello, Alfio Di Mauro, Andrea Borghesi, Andrea Cossettini, Angelo Garofalo, Arpan Prasad, Chi Zhang, Corrado Bonfanti, Cristian Cioflan, Cyril Koenig, Daniele Palossi, Davide Rossi, Fabio Montagna, Florian Glaser, Francesco Conti, Georg Rutishauser, Germain Haugou, Gianna Paulin, Giuseppe Tagliavini, Hanna Müller, Jannis Schoenleber, Lorenzo Lamberti, Luca Bertaccini, Luca Colagrande, Luca Valente, Maicol Caini, Manuel Eggimann, Manuele Rusci, Marco Bertuletti, Marco Guermandi, Matheus Cavalcante, Matteo Perotti, Mattia Sinigaglia, Michael Rogenmoser, Moritz Scherer, Moritz Schneider, Nazareno Bruschi, Nils Wistoff, Paul Scheffler, Philipp Mayer, Robert Balas, Samuel Riedel, Segio Mazzola, Sergei Vostrikov, Simone Benatti, Thomas Benz, Thorir Ingolfsson, Tim Fischer, Victor Javier Kartsch Morinigo, Victor Jung, Viviane Potocnik, Vlad Niculescu, Xiaying Wang, Yichao Zhang, Yvan Tortorella, Frank K. Gürkaynak, all our past collaborators
and many more that we forgot to mention



<http://pulp-platform.org>



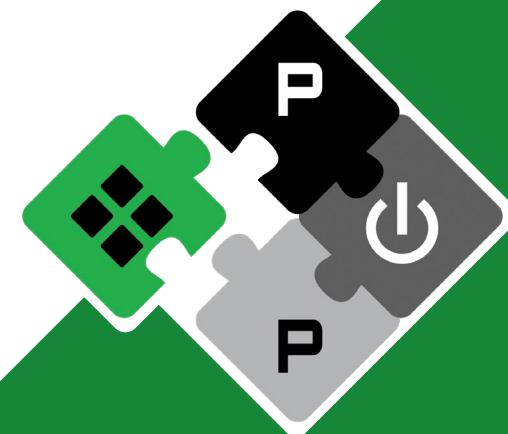
@pulp_platform



Additional Slides

PULP Platform

Open Source Hardware, the way it should be!

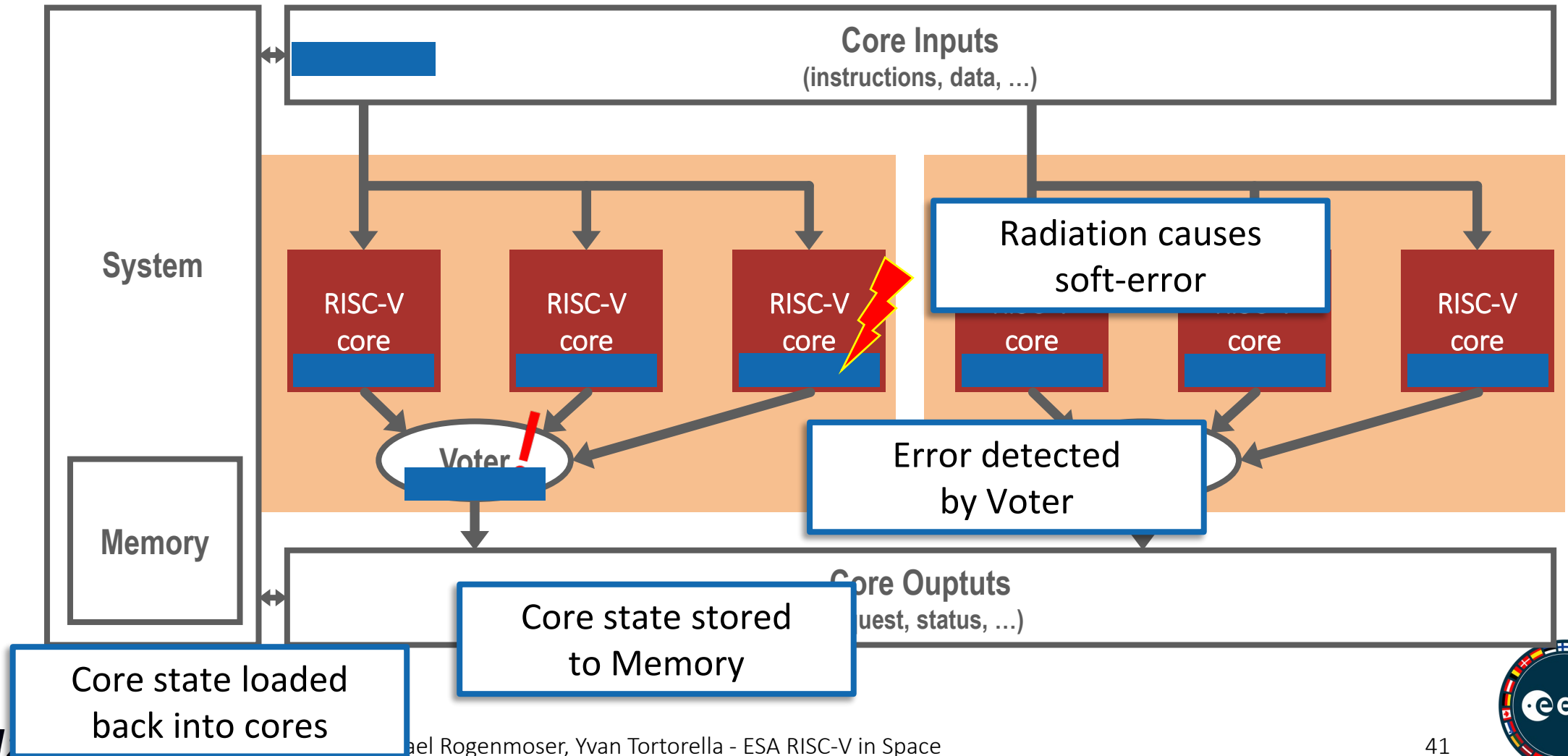


@pulp_platform 

pulp-platform.org 

youtube.com/pulp_platform 

Re-synchronization



RISC-V cores developed by PULP team



32 bit			64 bit
Low Cost Core	Core with DSP support	Core for Streaming	Linux capable Core
■ IBEX ■ Zero-riscy ■ Micro-riscy ■ 2 options ■ RV32-ICM ■ RV32-CE	■ CV32E40P ■ RV32-ICMXP ■ SIMD ■ Low loops ■ Bit manipulation	■ Snitch ■ RV32-IMAFD ■ Small Int core ■ Sign 64bit ■ FPU	■ Ariane/CVA6 ■ RV64-ICMAFD
<i>Maintained BY LowRISC</i> ARM Cortex-M0+	<i>Maintained BY OpenHW</i> ARM Cortex-M4F Fixed point	<i>Brand New Core</i> novel Extensions	<i>Maintained by OpenHW</i> ARM Cortex-A55

PULP is Cores+Interco+IO+Accelerators → Open Platform



RISC-V Cores

Ibex
RV32-ICM/CE

CV32E40P
RV32-ICMXF

Snitch
RV32-IMAFA

CVA6
RV64-ICMAFD

Peripherals

JTAG

SPI

UART

I2S

DMA

GPIO

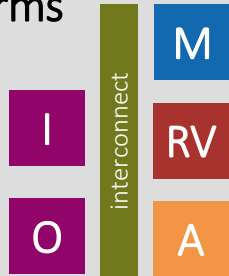
Interconnect

Logarithmic interconnect

APB – Peripheral Bus

AXI4 – Interconnect

Platforms



Single Core

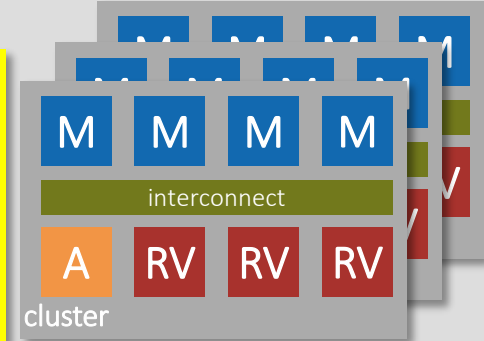
- PULPissimo

Open-Source Hardware > Open-Source ISA:

- core implementations
- system-level integration
- peripherals, accelerators

Everything
designed in
System Verilog

- Mr. Wolf
- Vega



Multi-cluster

- HERO
- Open Piton

IOT

HPC

Hardware Accelerators

NE
(DNN)

NTX
(ML)

FFT

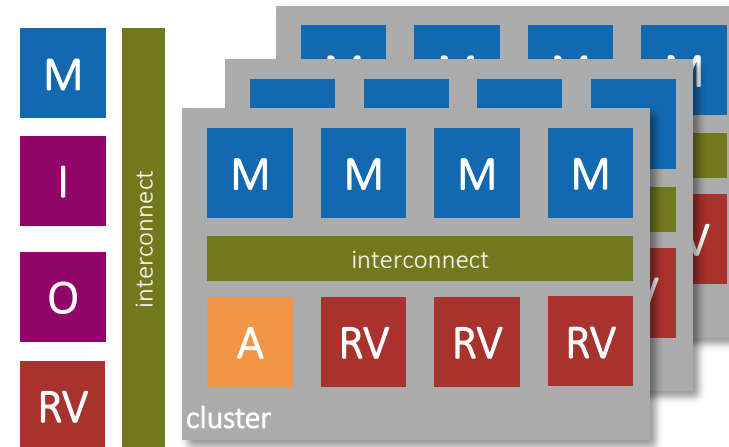
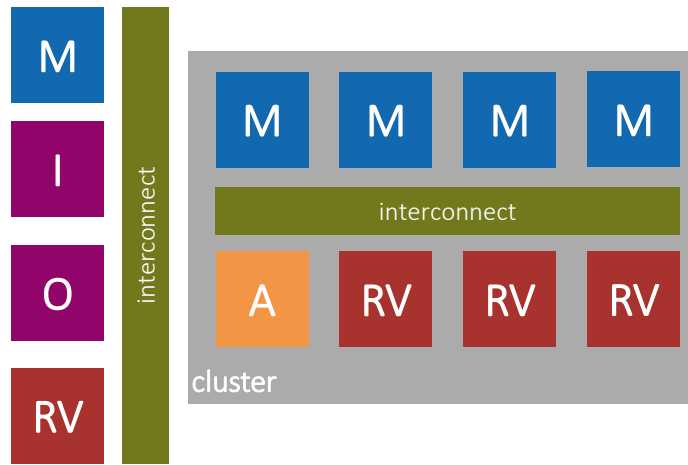
RedMule
(MatMul)

Heterogeneity of the PULP Ecosystem

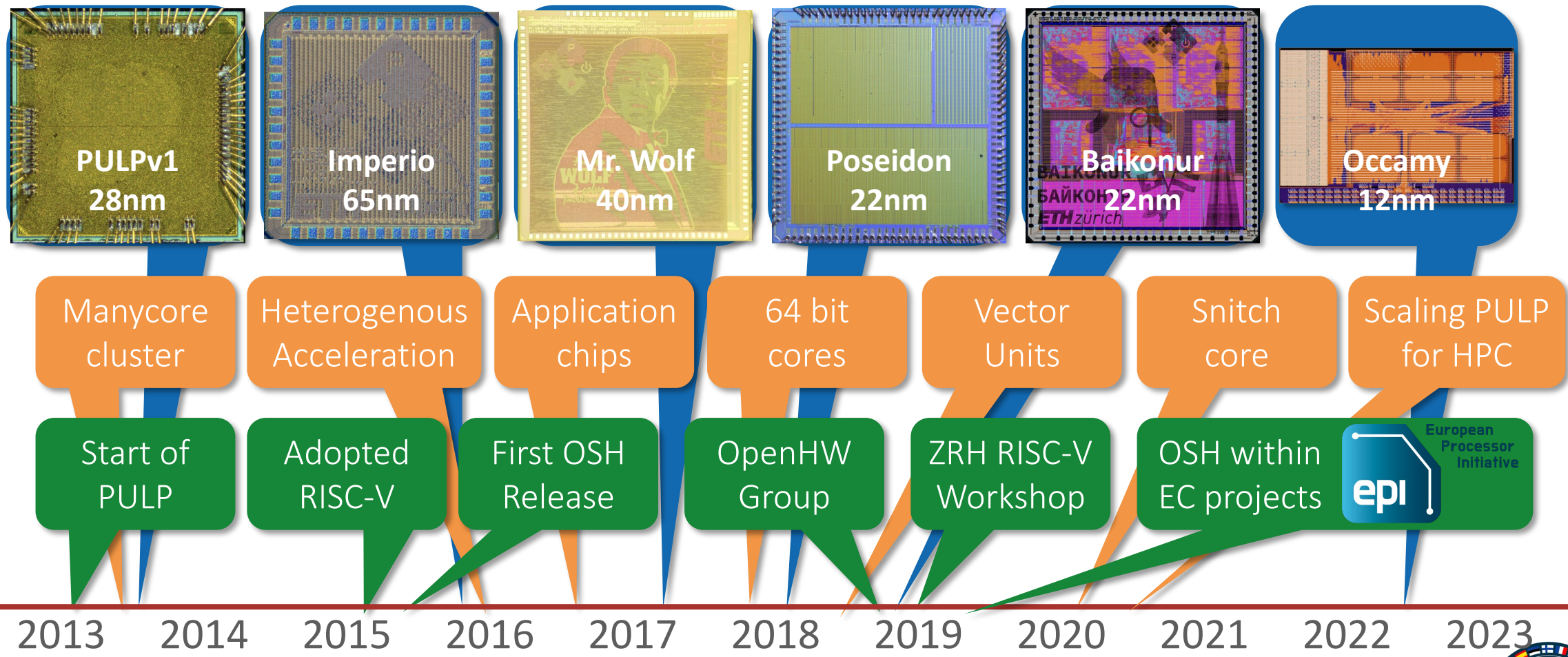


Heterogeneity in the PULP Ecosystem can be exploited at multiple levels

- **Cluster Level:** parallel **general-purpose** RISC-V cores + **application-specific** hardware accelerators
- **Address Space Level:** **64-bit** controller core + **32-bit** general-purpose multi-core (or multi-cluster) accelerator

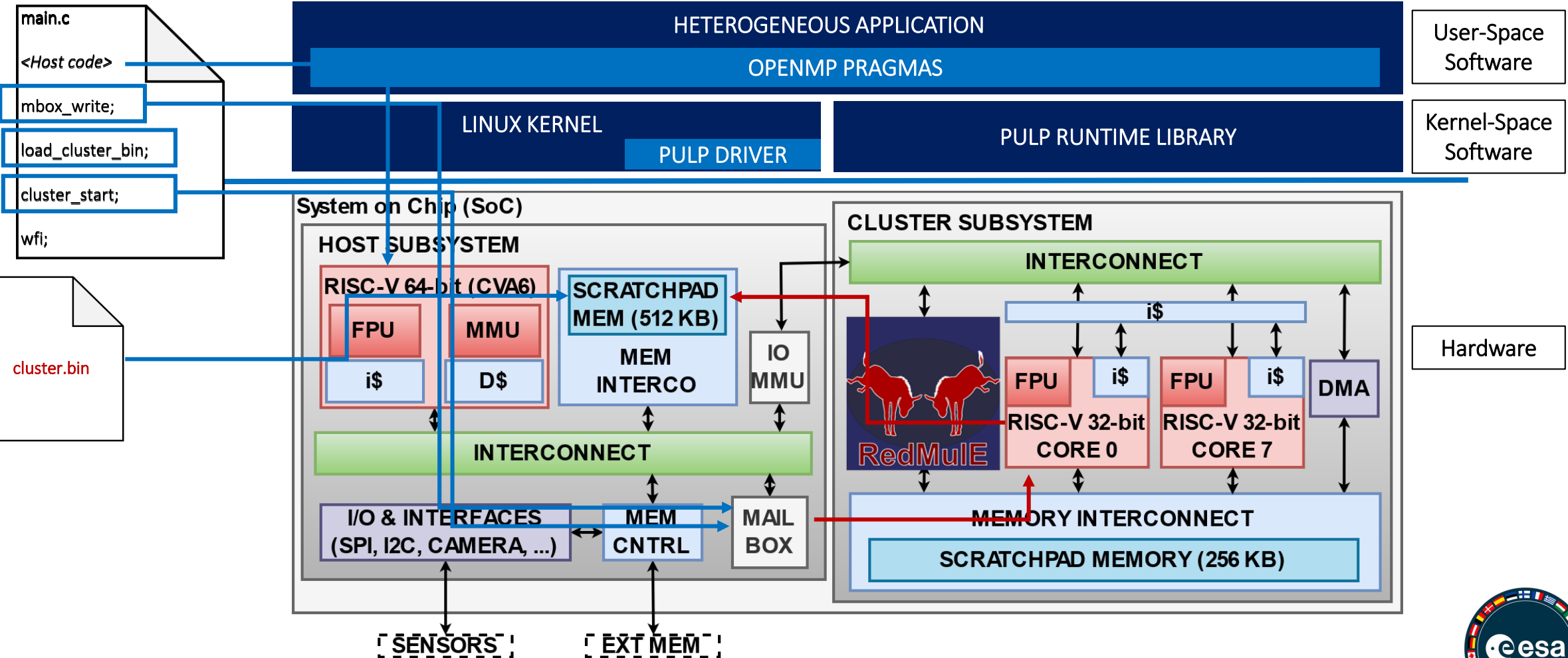


Timeline of Parallel Ultra Low Power (PULP) project



Heterogeneous SoC for Space Applications

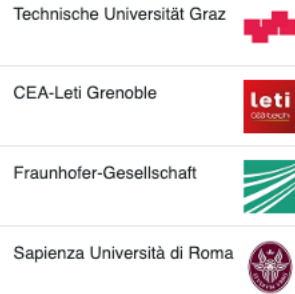
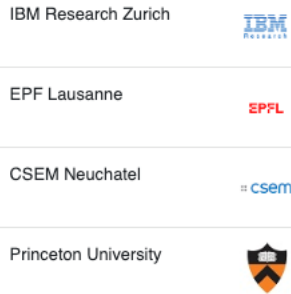
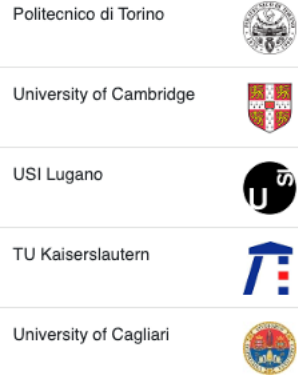
Bare-through-Linux Offloading



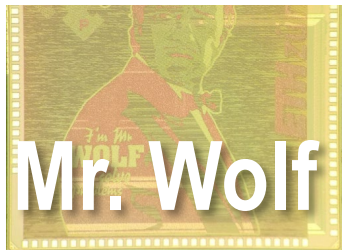
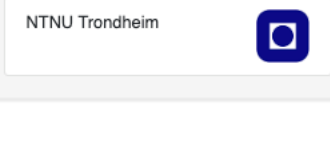
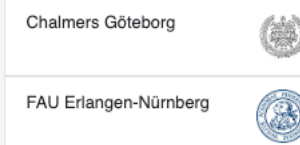
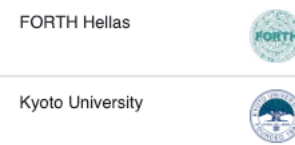
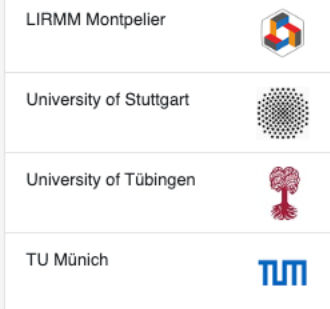
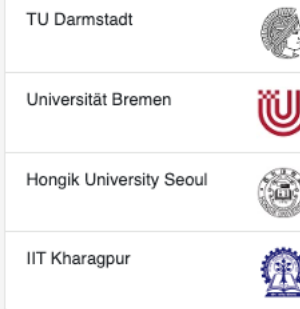
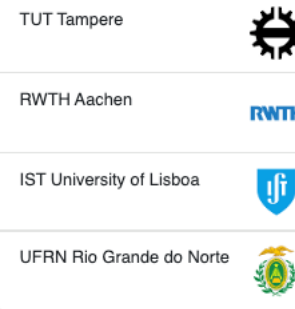
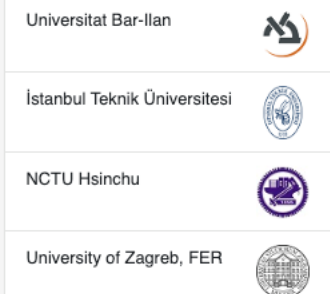
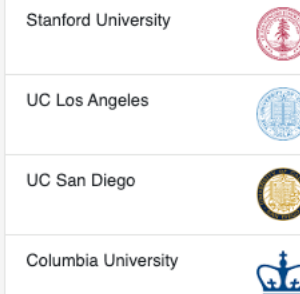
Open source collaboration scheme explained



Direct research collaborators on PULP



Academic users we are aware of



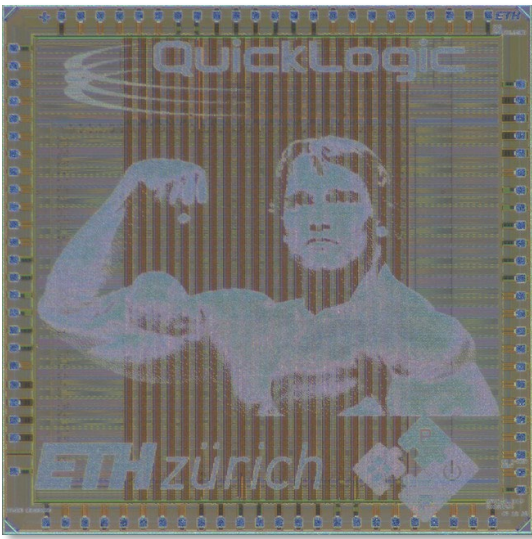
commits

commits

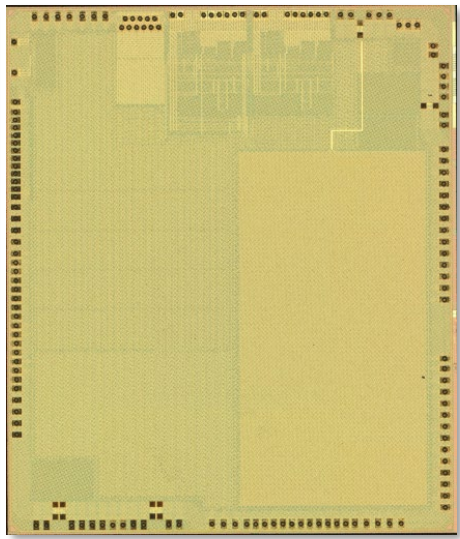
GitHub



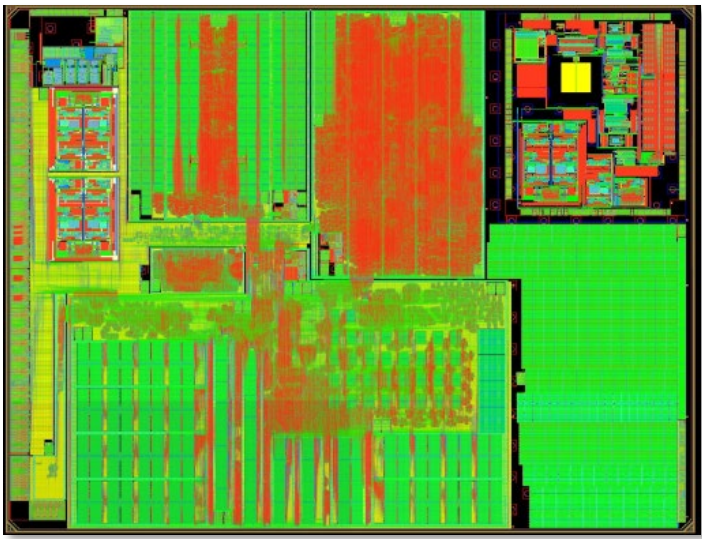
The open model led to successful industry collaborations



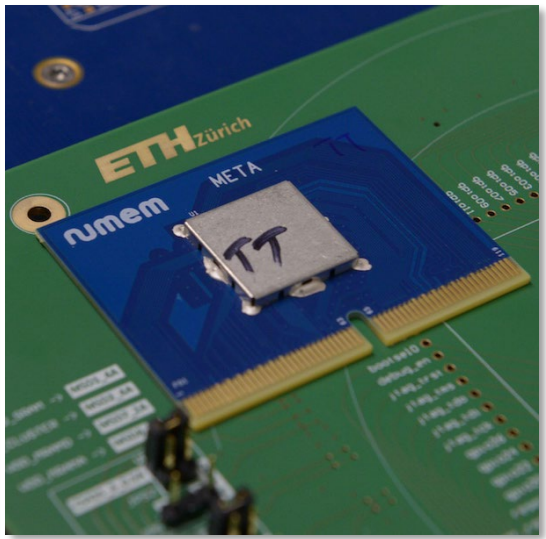
Arnold (GF22)
eFPGA with RISC-V core



Vega (GF22)
IoT Processor with ML acceleration



Marsellus (GF22)
IoT Processor with low power modes and event based computing



Siracusa (TSMC16)
IoT Processor with NVM technology

