



PULP PLATFORM

Open Source Hardware, the way it should be!

Open, Parallel Ultra-Low Power Platforms for Extreme Edge AI

Luca Benini <lbenini@iis.ee.ethz.ch>



European
Commission

Horizon 2020
European Union funding
for Research & Innovation



FONDS NATIONAL SUISSE
SCHWEIZERISCHER NATIONALFONDS
FONDO NAZIONALE SVIZZERO
SWISS NATIONAL SCIENCE FOUNDATION



<http://pulp-platform.org>

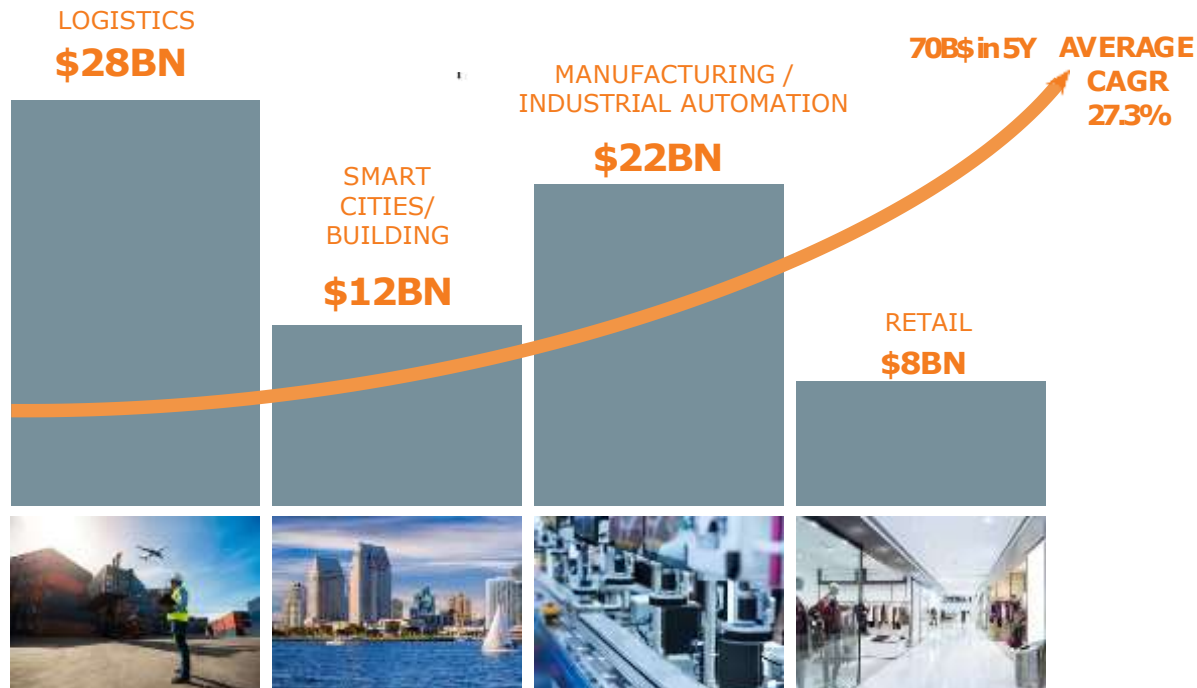
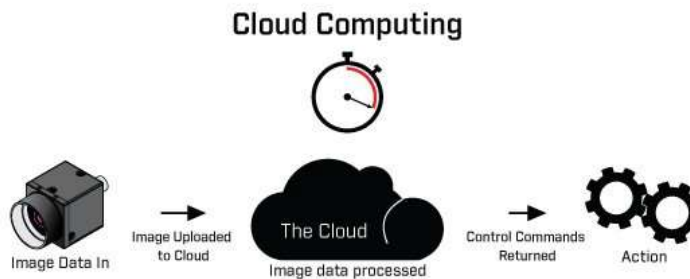


@pulp_platform

ETH zürich



Cloud → Edge → Extreme Edge AI a.k.a. TinyML



#1 Customer Question on Amazon.com (out of 1,000+):

1. I don't want any of my (private, personal) videos on any servers not in my control. Is this possible?

Source: www.amazon.com/ask/questions/asins/B01N5VH087/

#2 Customer Question on Amazon.com (out of 1,000+):

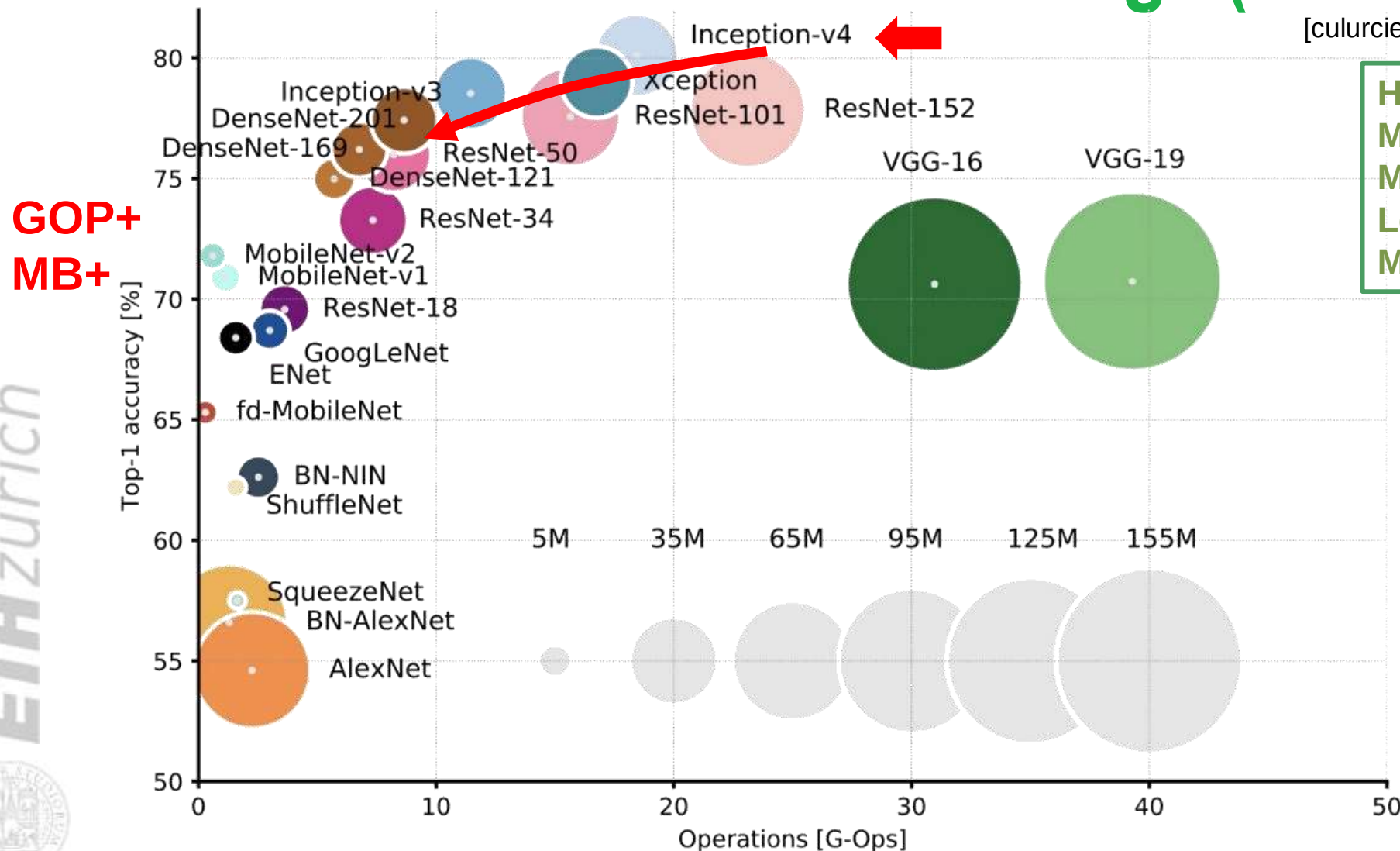
2. How long does the battery charge last?

Source: www.amazon.com/ask/questions/asins/B01N5VH087/

Extreme edge AI challenge
AI capabilities in the power envelope of an MCU:

100mW peak (1mW avg)

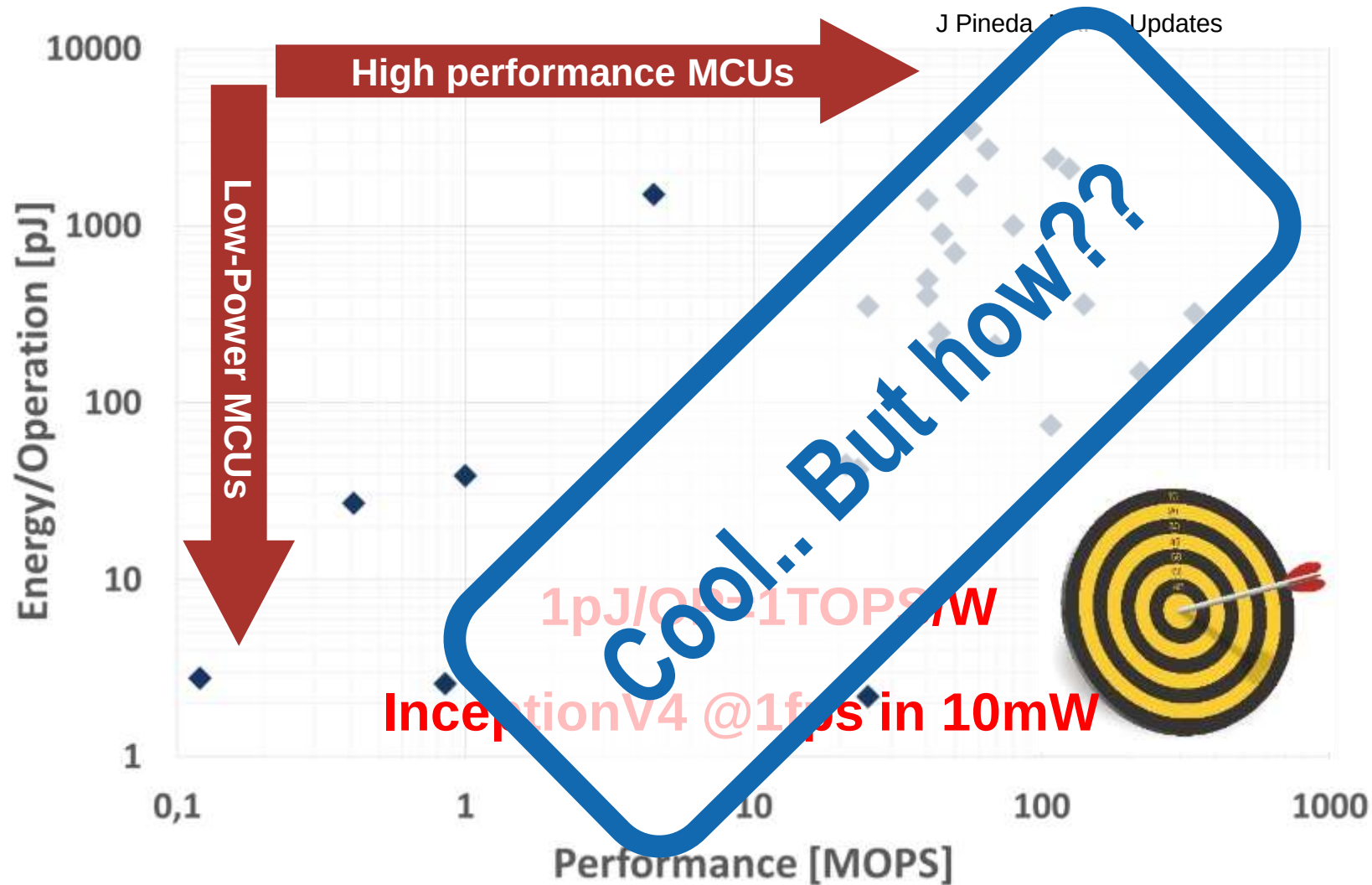
AI Workloads from Cloud to Edge (Extreme?)



[culurciello18]

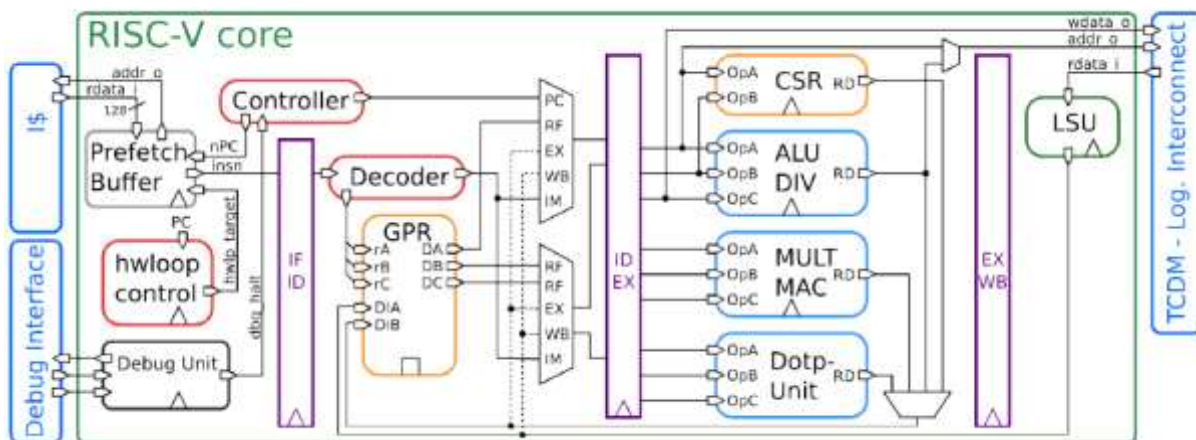
High OP/B ratio
Massive Parallelism
MAC-dominated
Low precision OK
Model redundancy

Energy efficiency @ GOPS is **THE** Challenge

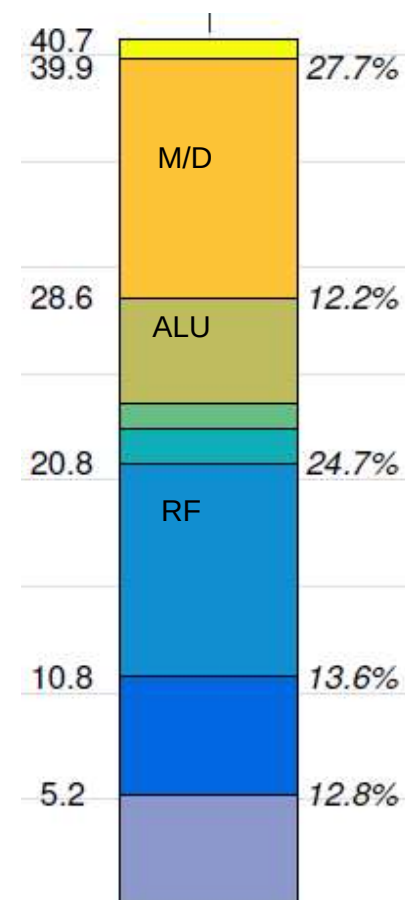


RI5CY – An Open MCU-class RISC-V Core for EE-AI

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



40 kGE
70% RF+DP



RISC-V ISA is extensible *by construction* (great!)

V1 Baseline RISC-V RV32IMC (not good for ML)

HW loops

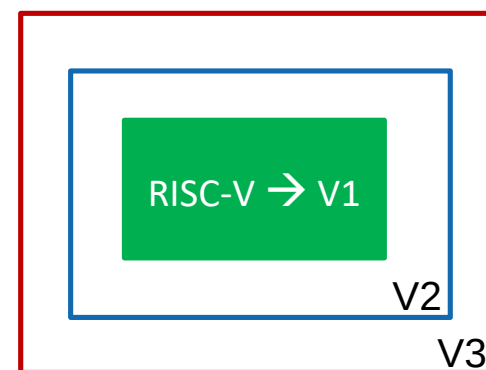
V2 Post modified Load/Store

Mac

V3 SIMD 2/4 + DotProduct + Shuffling

Bit manipulation unit

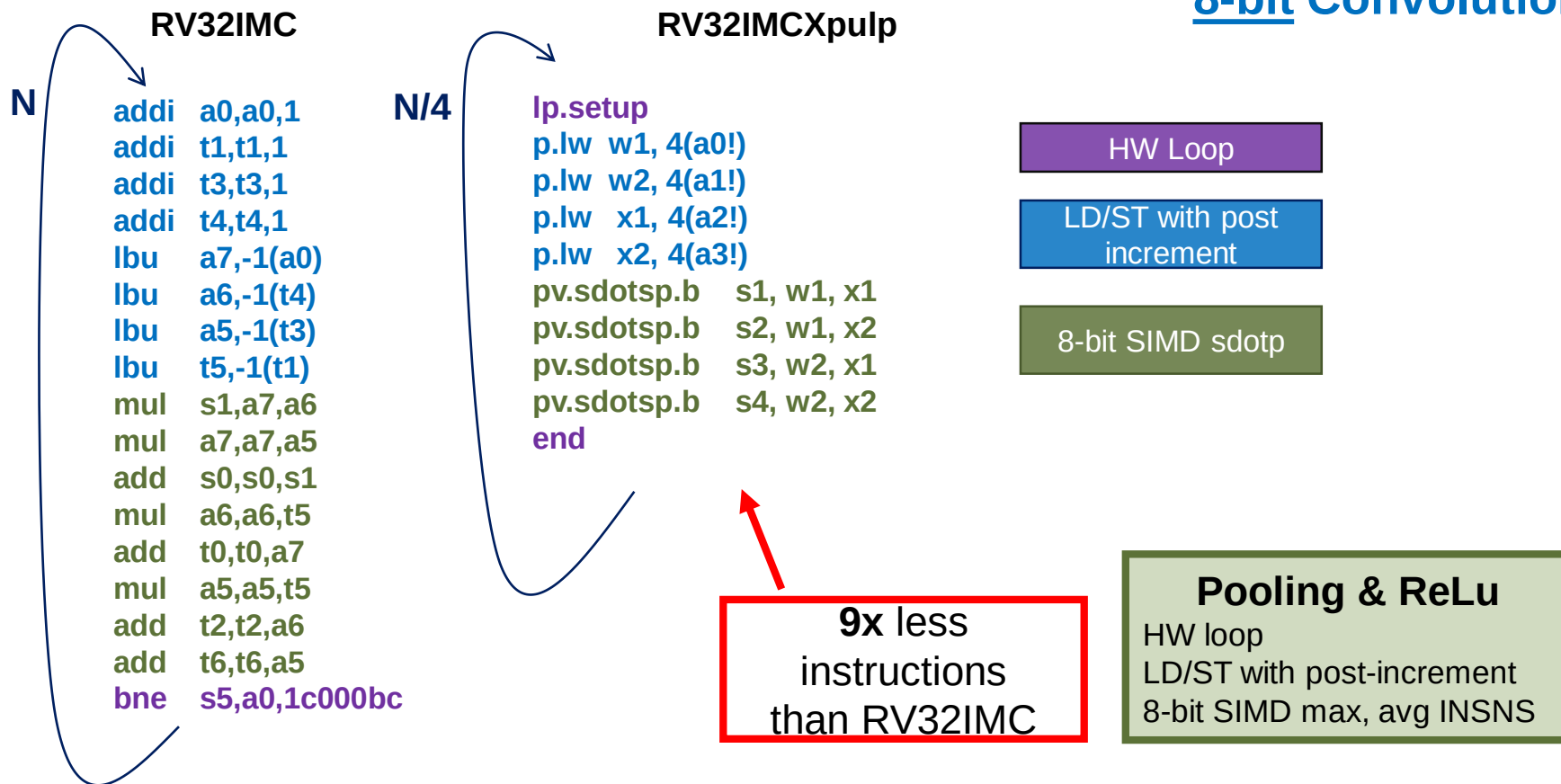
Lightweight fixed point



XPULP extensions: 25 kGE → 40 kGE (1.6x)

PULP-NN: Xpulp ISA exploitation

8-bit Convolution



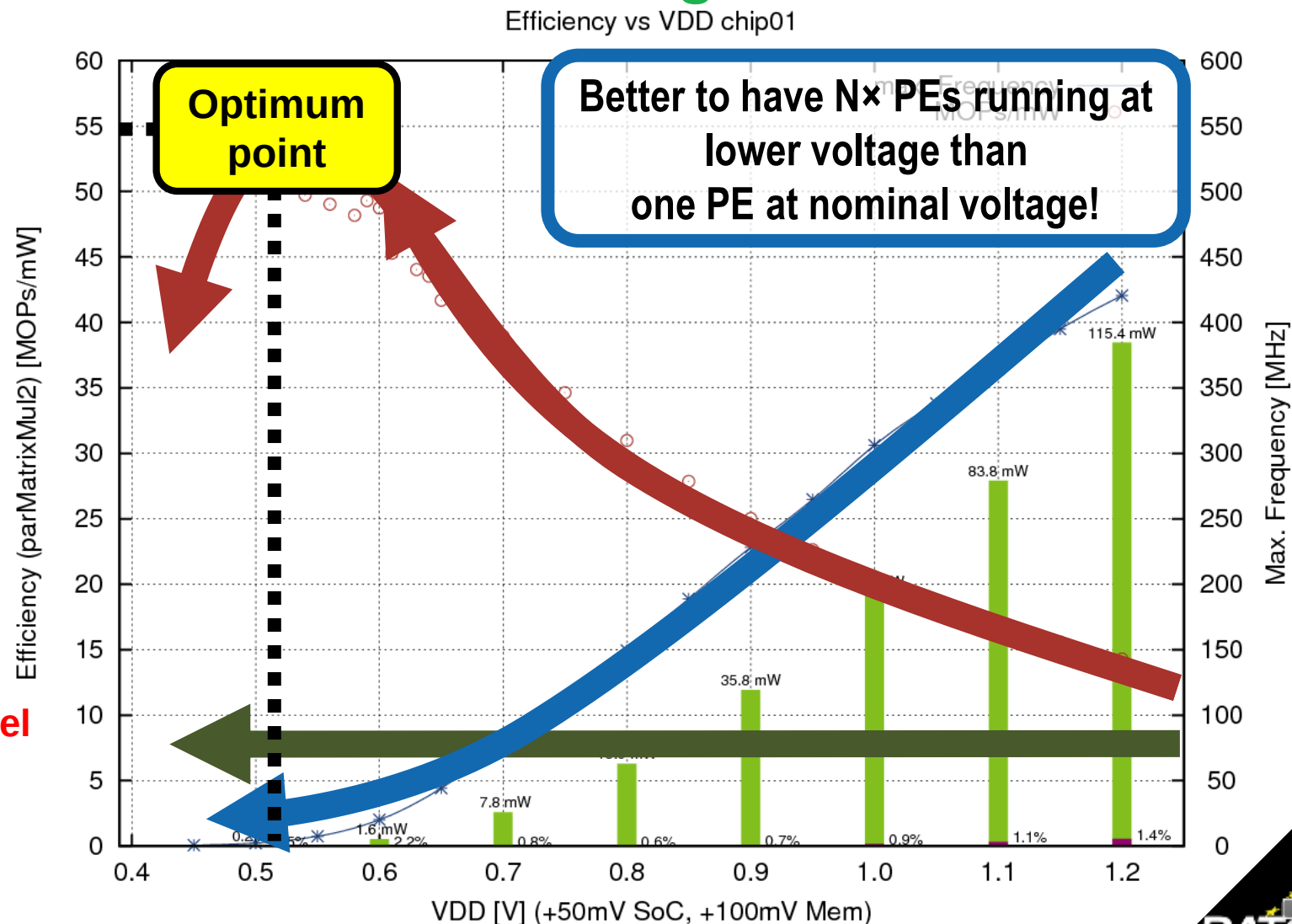
Nice – But what about the GOPS?
Faster+Superscalar is not efficient!

M7: 5.01 CoreMark/MHz-58.5 μ W/MHz
 M4: 3.42 CoreMark/MHz-12.26 μ W/MHz

ML & Parallel, Near-threshold: a Marriage Made in Heaven

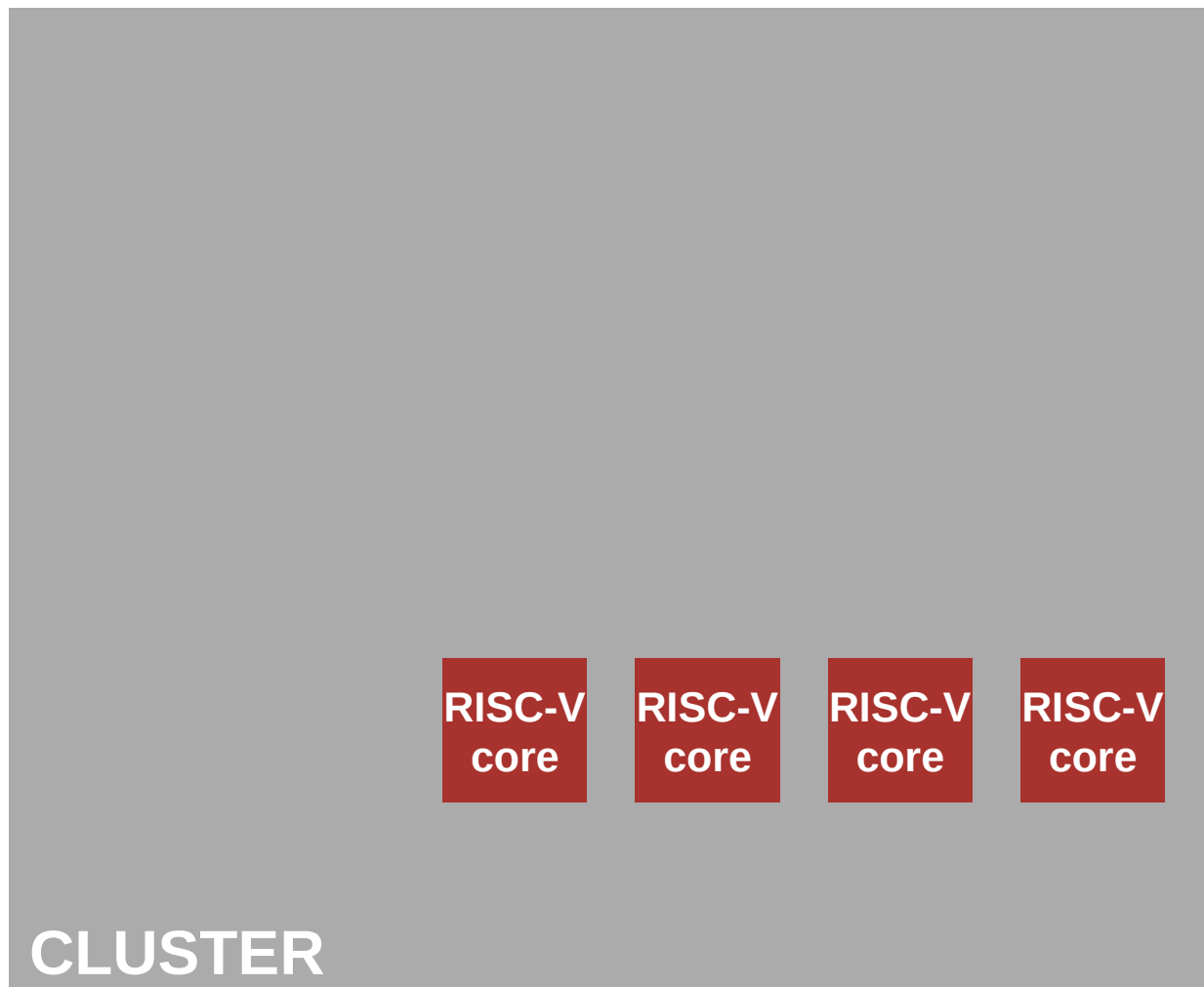
- As **VDD** decreases, **operating speed** decreases
- However **efficiency** increases → more work done per Joule
- Until leakage effects start to dominate
- Put more units in parallel to get performance up and keep them busy with a parallel workload

ML is massively parallel and scales well (P/S ↑ with NN size)





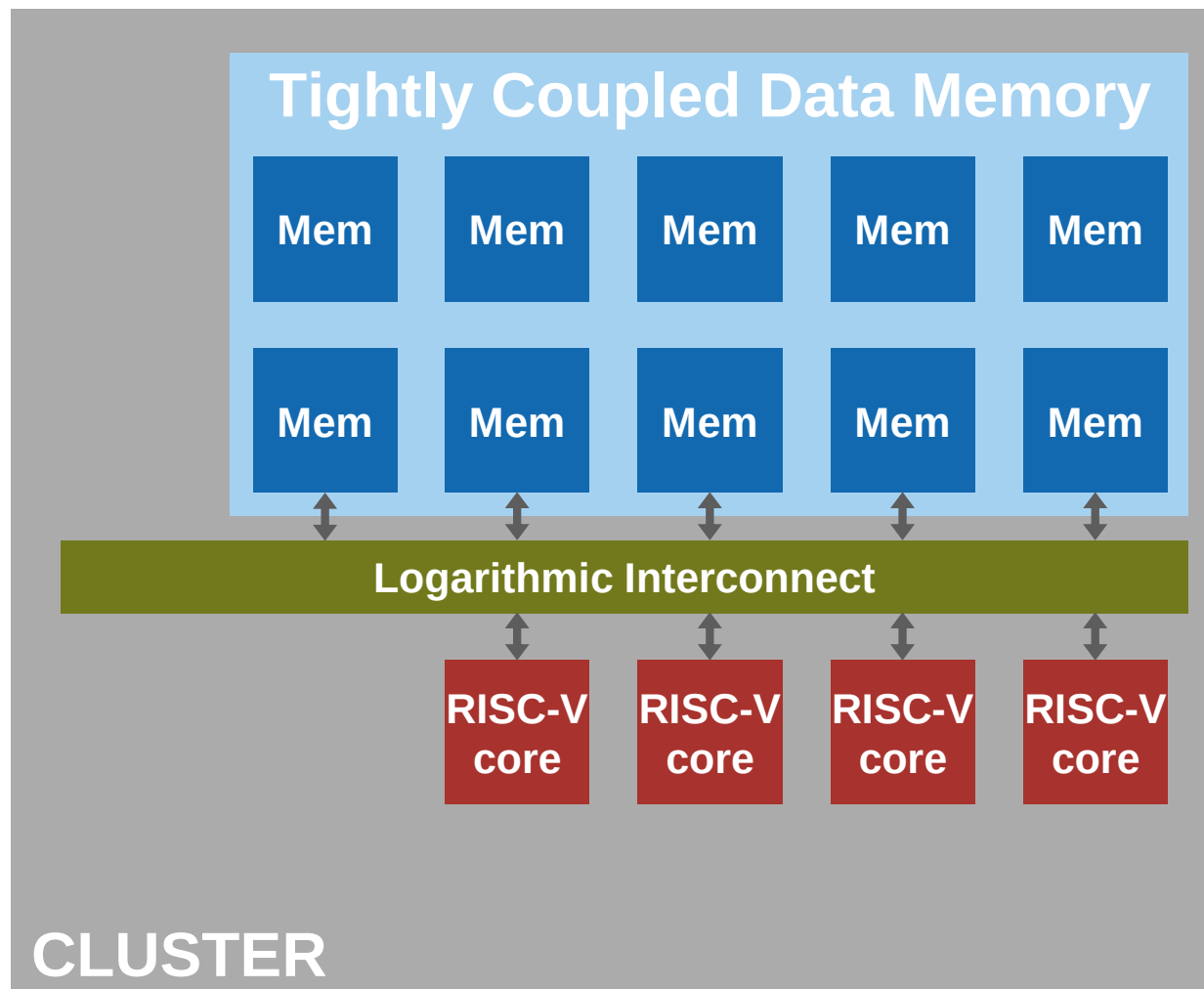
Multiple RI5CY Cores (1-16)



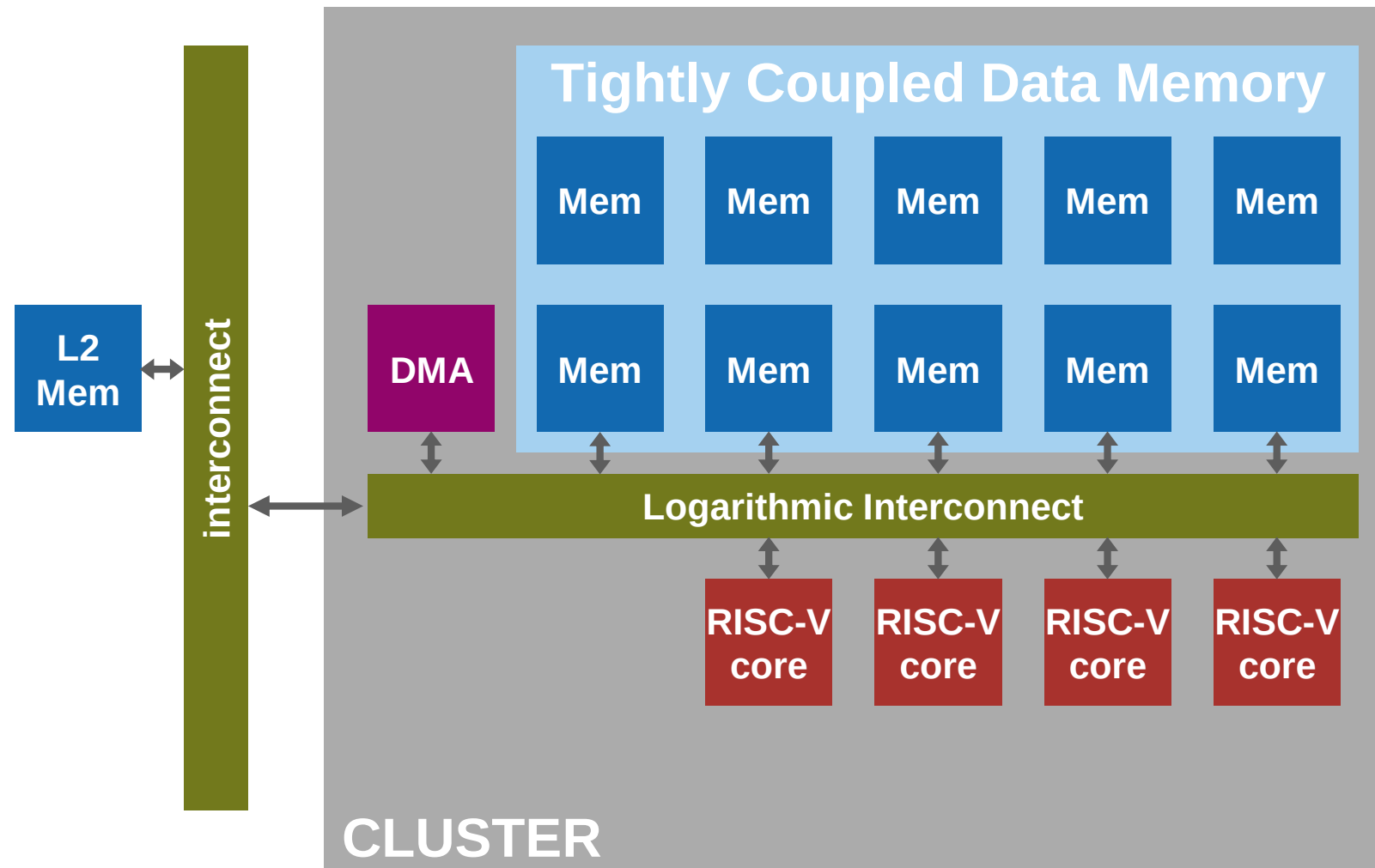
ETH zürich



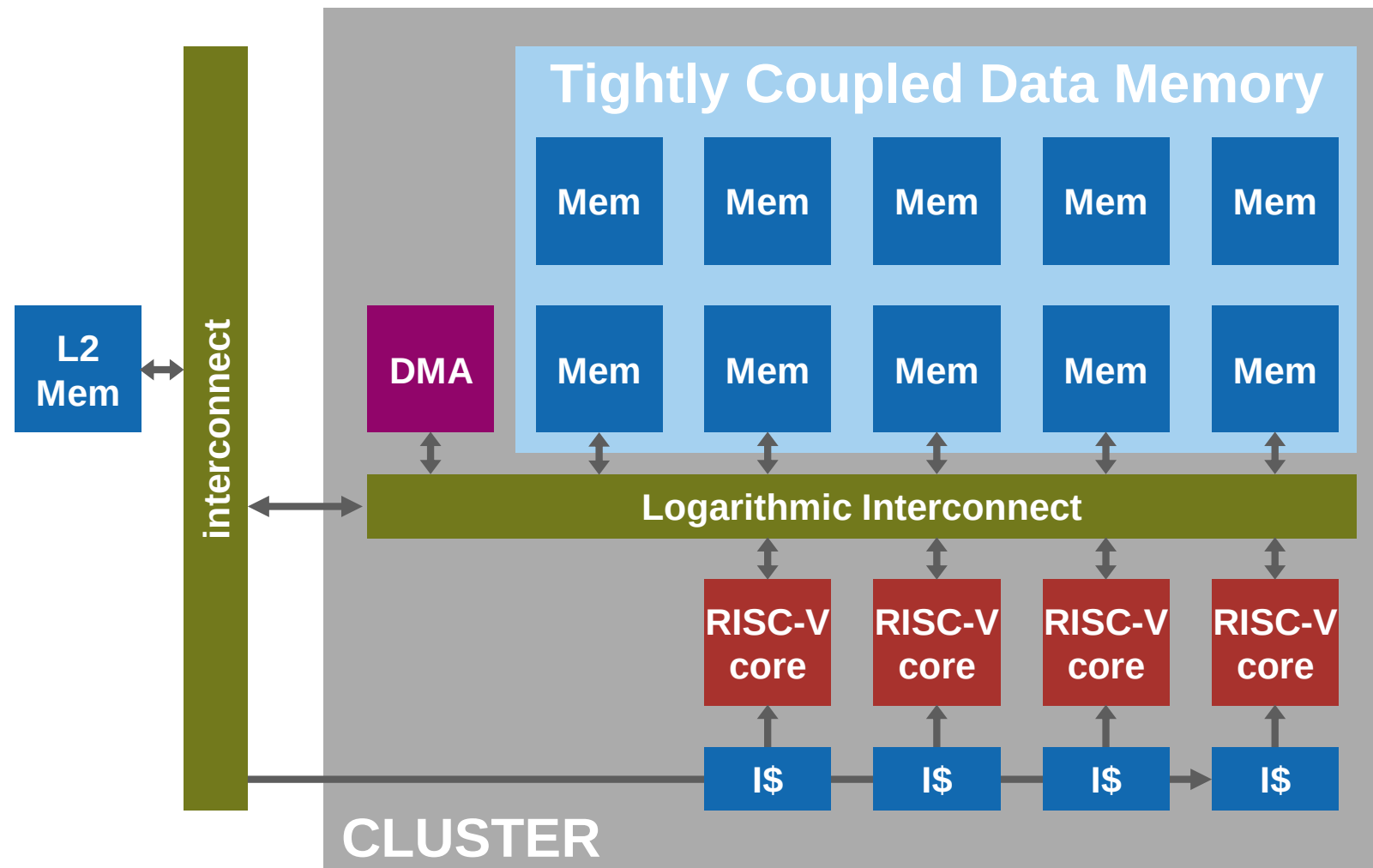
Low-Latency Shared TCDM



DMA for data transfers from/to L2



Shared instruction cache with private “loop buffer”



Results: RV32IMCxpulp vs RV32IMC

■ 8-bit convolution

- Open source DNN library

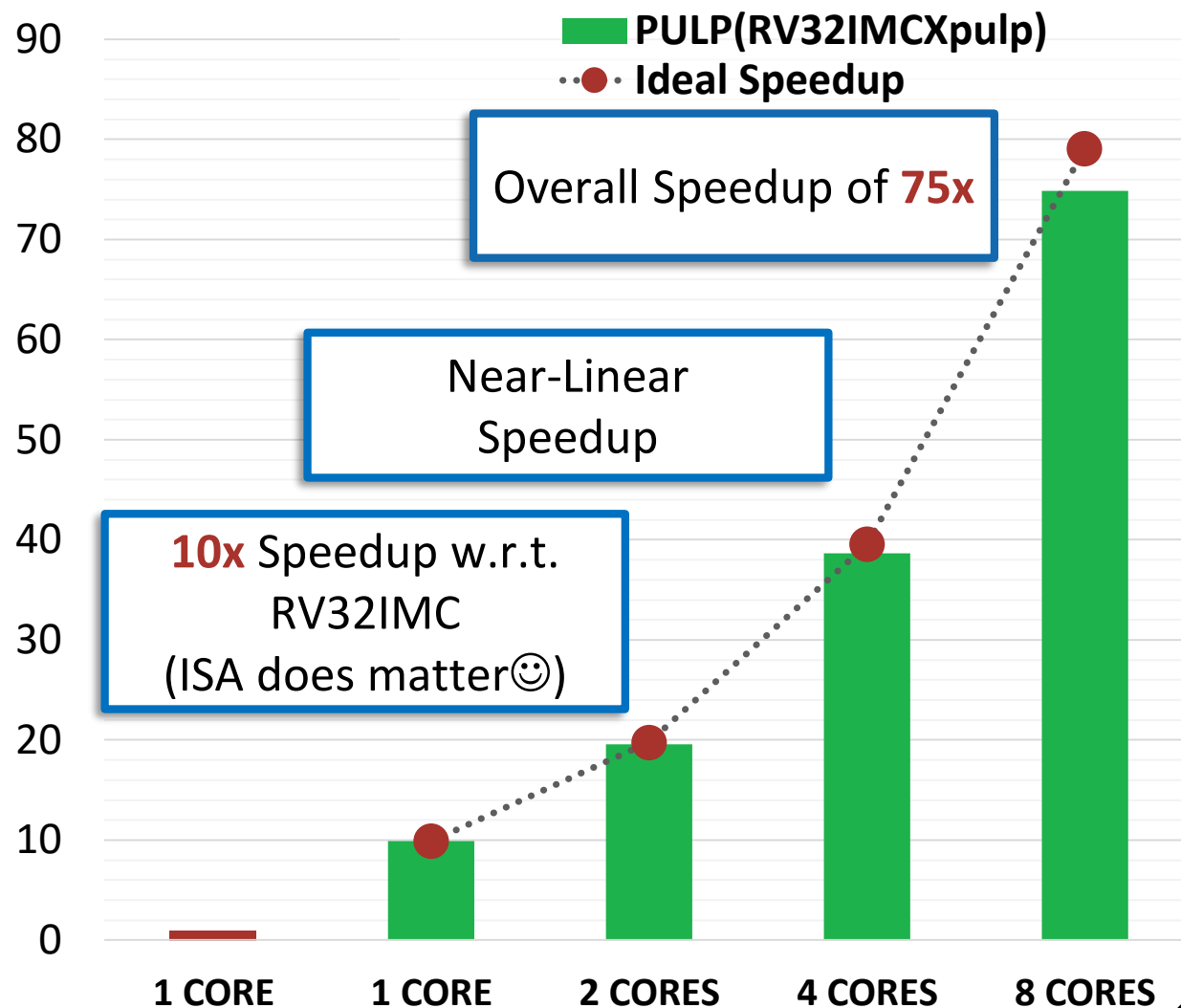
■ 10x through xPULP

- Extensions bring real speedup

■ Near-linear speedup

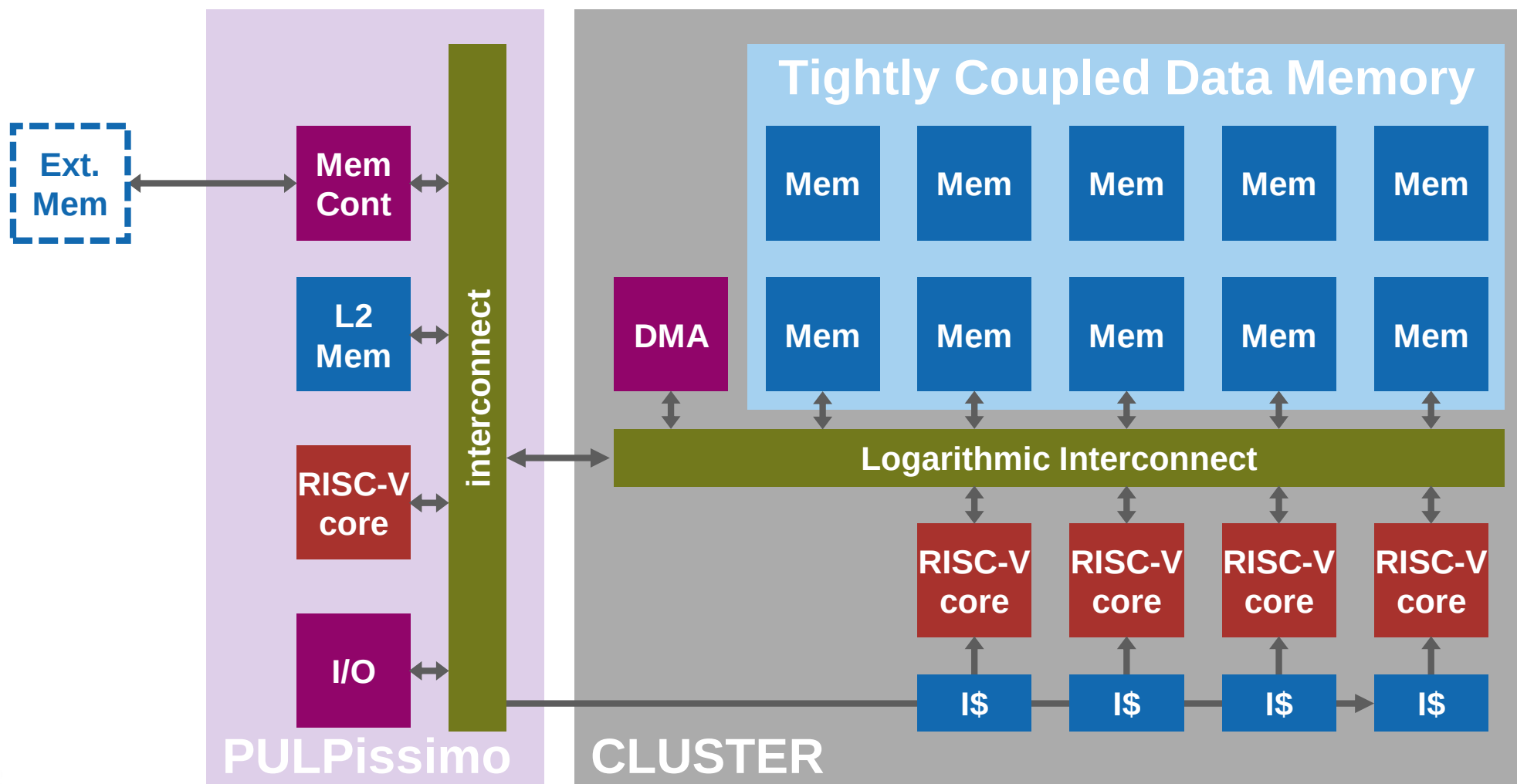
- Scales well for regular workloads.

■ 75x overall gain



PULP-NN: an open Source library for DNN inference on PULP cores

An additional I/O controller is used for IO



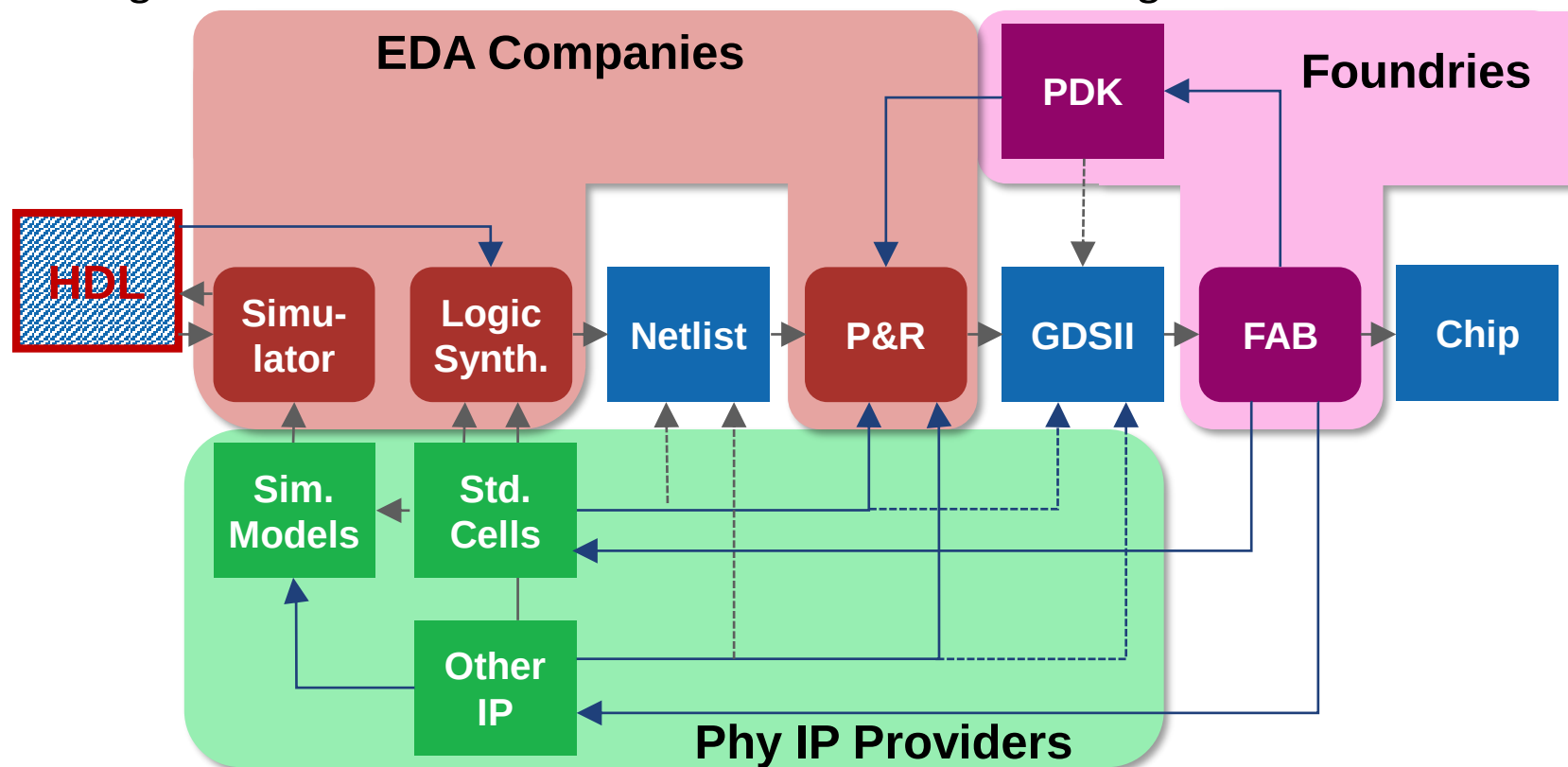
All this is Open Source HW – PULP open

Nice, but what exactly is “open” in Open Source HW?



Nice, but what exactly is “open” in Open Source HW?

- Only the first stage of the silicon production pipeline can be open HW
→ **RTL source code** (in an HDL such as SystemVerilog)
- Later stages contain closed IP of various actors + tool licensing issues



Permissive, Copyright (e.g APACHE) License is Key for industrial adoption

PULP includes Cores+Interco+IO → and Open Platform

RISC-V Cores

RI5CY

32b

Ibex

32b

Snitch

32b

Ariane
+ Ara
64b

Peripherals

JTAG

UART

DMA

SPI

I2S

GPIO

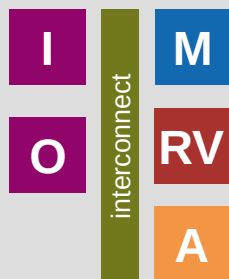
Interconnect

Logarithmic interconnect

APB – Peripheral Bus

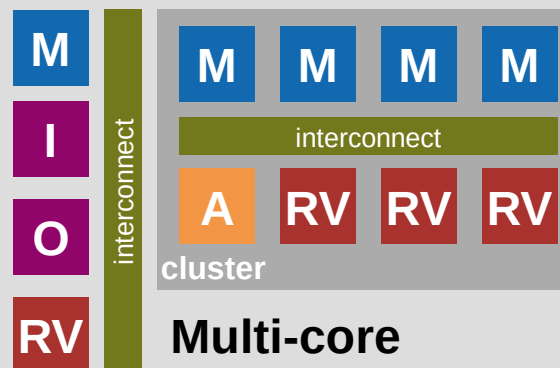
AXI4 – Interconnect

Platforms



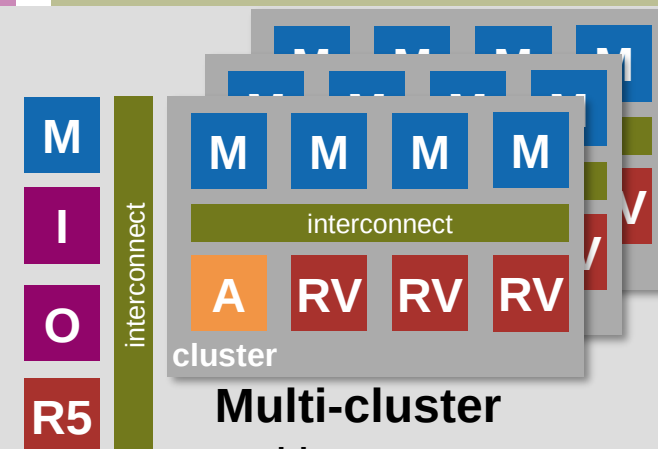
Single Core

- PULPino
- PULPissimo



Multi-core

- Fulmine
- Mr. Wolf



Multi-cluster

- Hero
- Open Piton

IOT

HPC

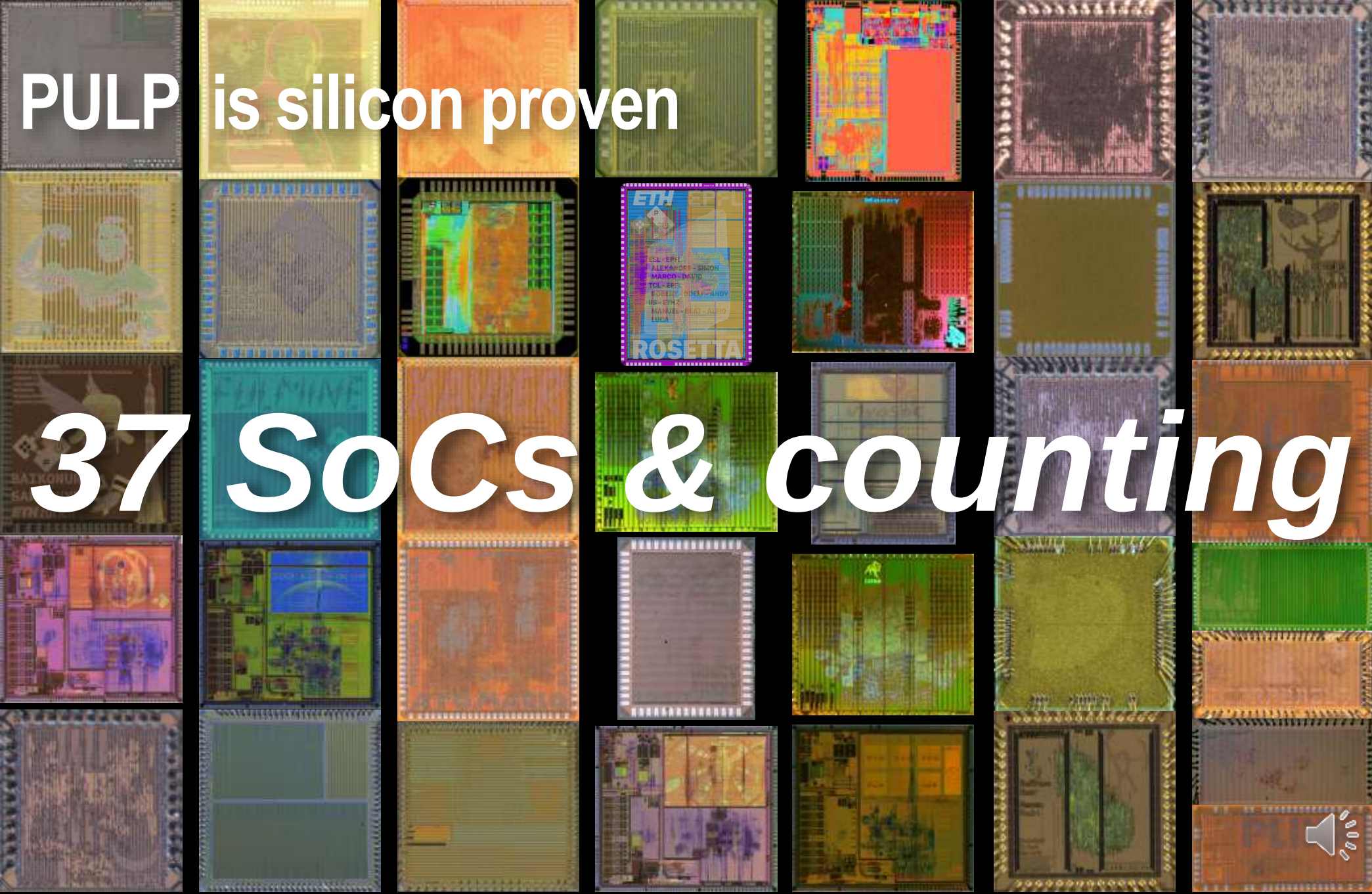
Accelerators

HWCE
(convolution)Neurostream
(ML)HWCrypt
(crypto)PULPO
(1st ord. opt)

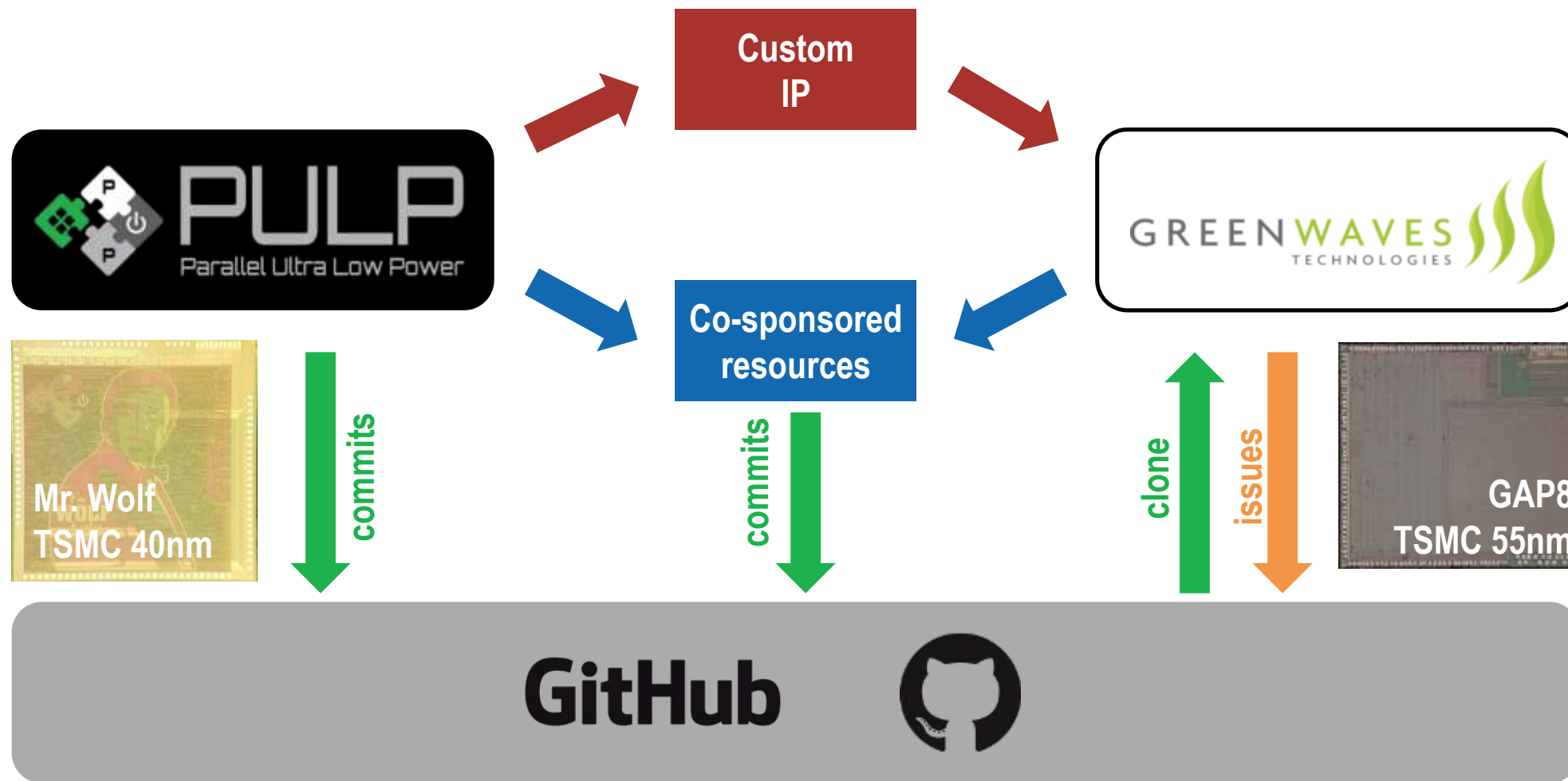
<http://asic.ethz.ch>

PULP is silicon proven

37 SoCs & counting

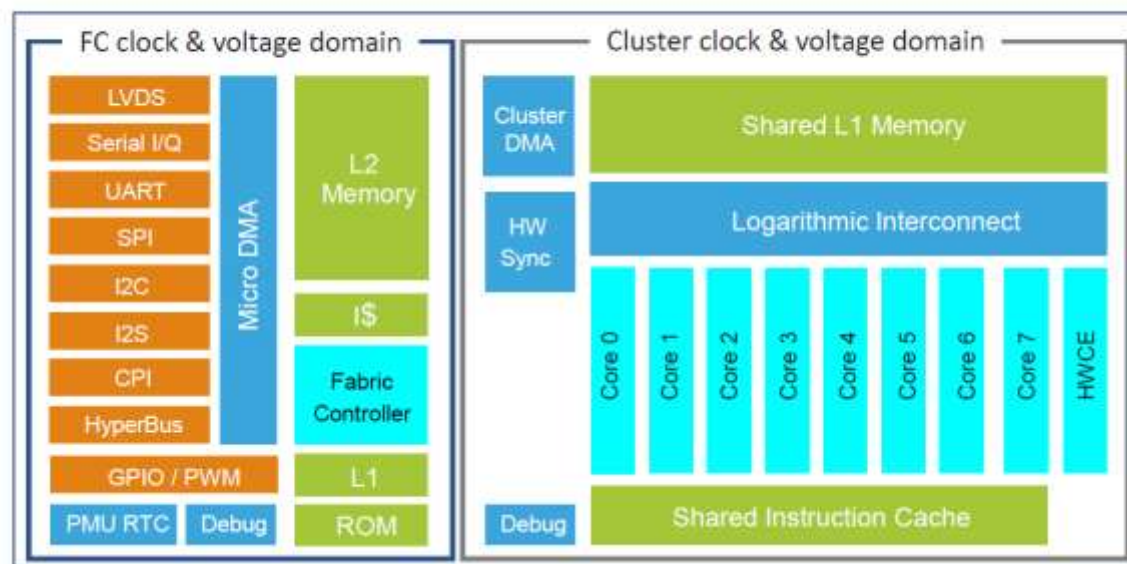


Open source collaboration with industry



Successful product development: GWT's GAP8

Two independent clock and voltage domains, from 0-133MHz/1V up to 0-250MHz/1.2V



MCU Function

- Extended RISC-V core
- Extensive I/O set
- Micro DMA
- Embedded DC/DC converter
- Secured execution

Computation engine

- 8 extended RISC-V cores
- Fully programmable
- Efficient parallelization
- Shared instruction cache
- Multi channel DMA
- HW synchronization
- HW convolution Engine



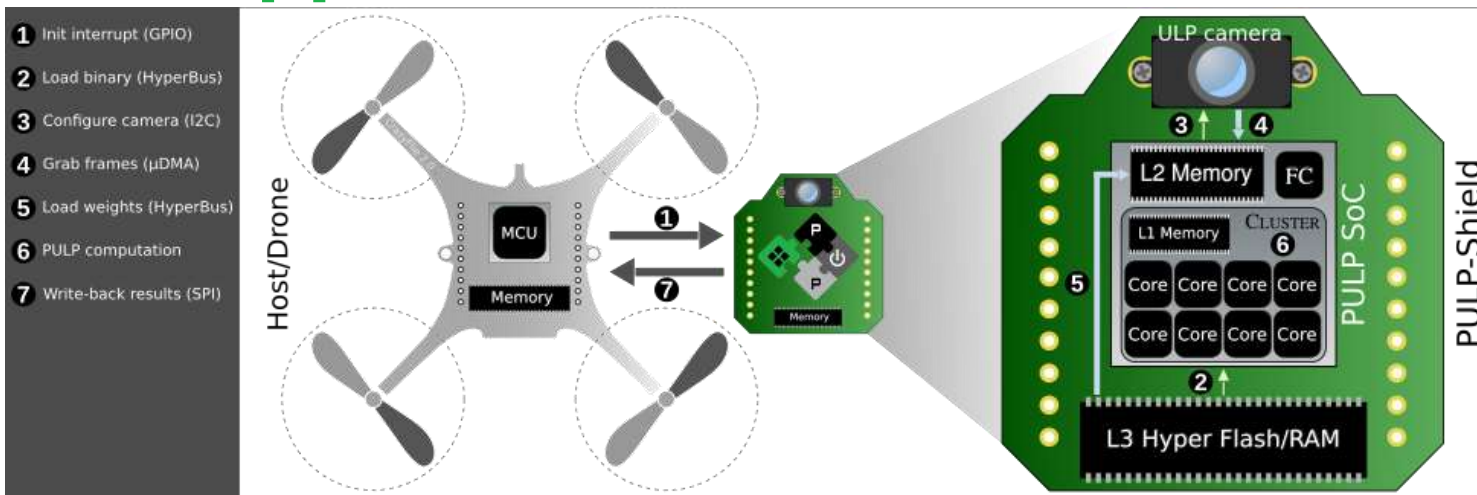
What	Freq MHz	Exec Time ms	Cycles	Power mW
40nm Dual Issue MCU	216	99.1	21 400 000	60
GAP8 @1.0V	15.4	99.1	1 500 000	3.7
GAP8 @1.2V	17.5	8.7	1 500 000	70
GAP8 @1.0V w HWCE	4.7	99.1	460 000	0.8

11 X

16 X

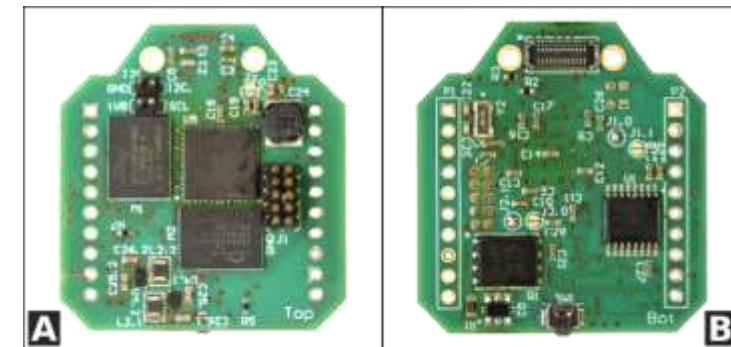


New Application Frontiers: DroNET on NanoDrone



Pluggable PCB:
PULP-Shield

- ~5g, 30×28mm
- GAP8 SoC
- 8 MB HDRAM
- 16 MB HFlash
- QVGA ULP HiMax camera
- Crazyflie 2.0 nano-drone (27g)



**Only onboard computation for autonomous flight + obstacle avoidance
no human operator, no ad-hoc external signals, and no remote base-station!**

Outlook: PULP-based nano UAV family



Copyright 2020 © **ETH** zürich

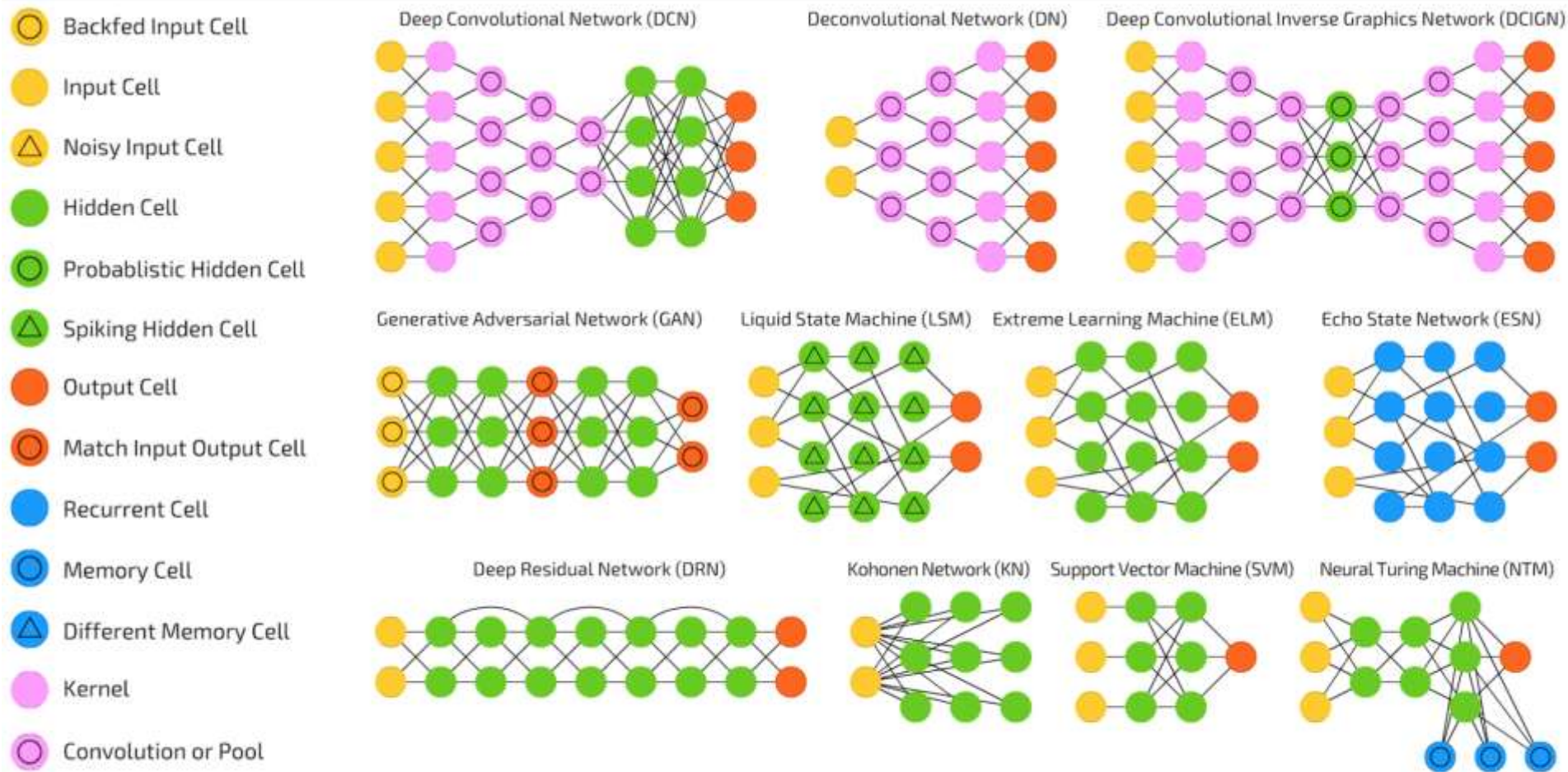
Credit: Hanna Müller & Daniele Palossi

- FrontNET [1]: hovering in front of a freely moving user
- End-to-end, RESnet-like CNN
- Runs with DroneNET
- AI+control: 5-10% of tot. Power
- Collaboration with IDSIA (USI-SUPSI) Lugano, Switzerland



[1] Mantegazza, Dario, et al. "Vision-based Control of a Quadrotor in User Proximity: Mediated vs End-to-End Learning Approaches." 2019 International Conference on Robotics and Automation (ICRA). IEEE, 2019.

What's next? Sub-pJ/OP Accelerators, but flexibility needed!

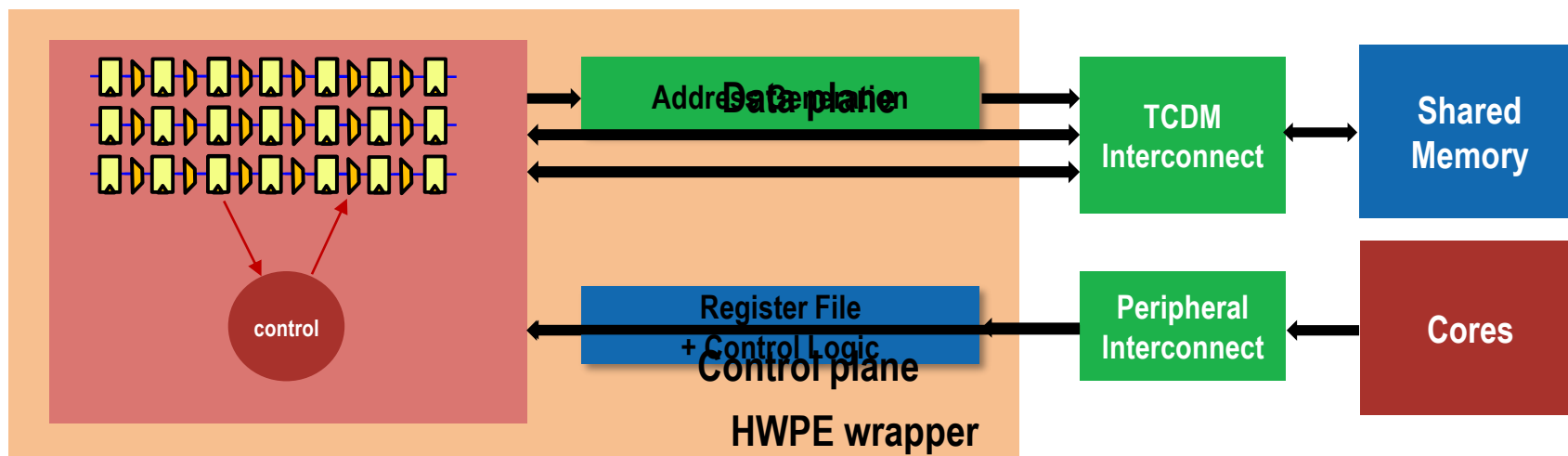
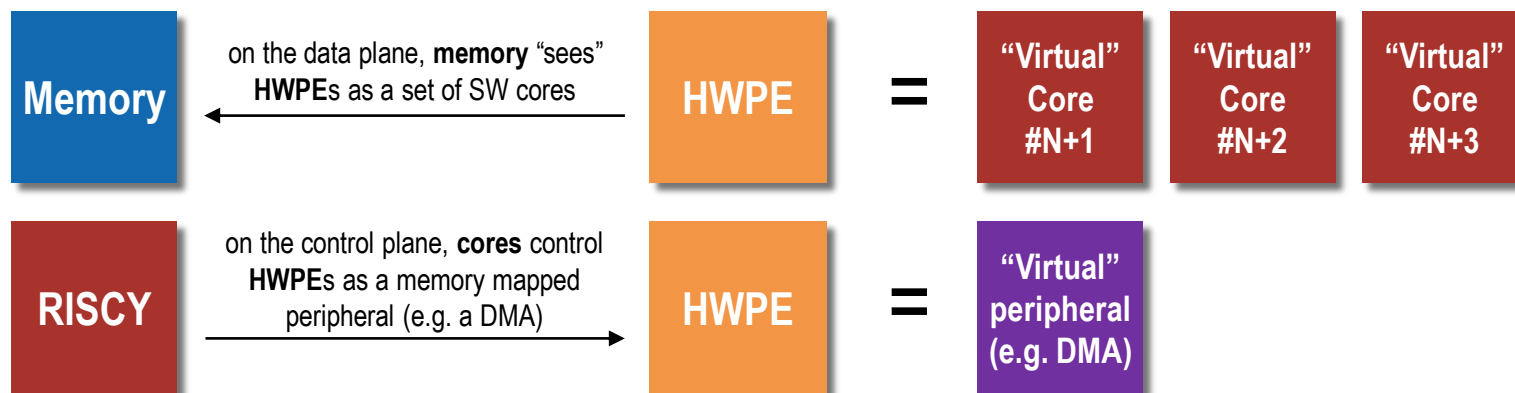


+ FFT, PCA, Mat-inv,...

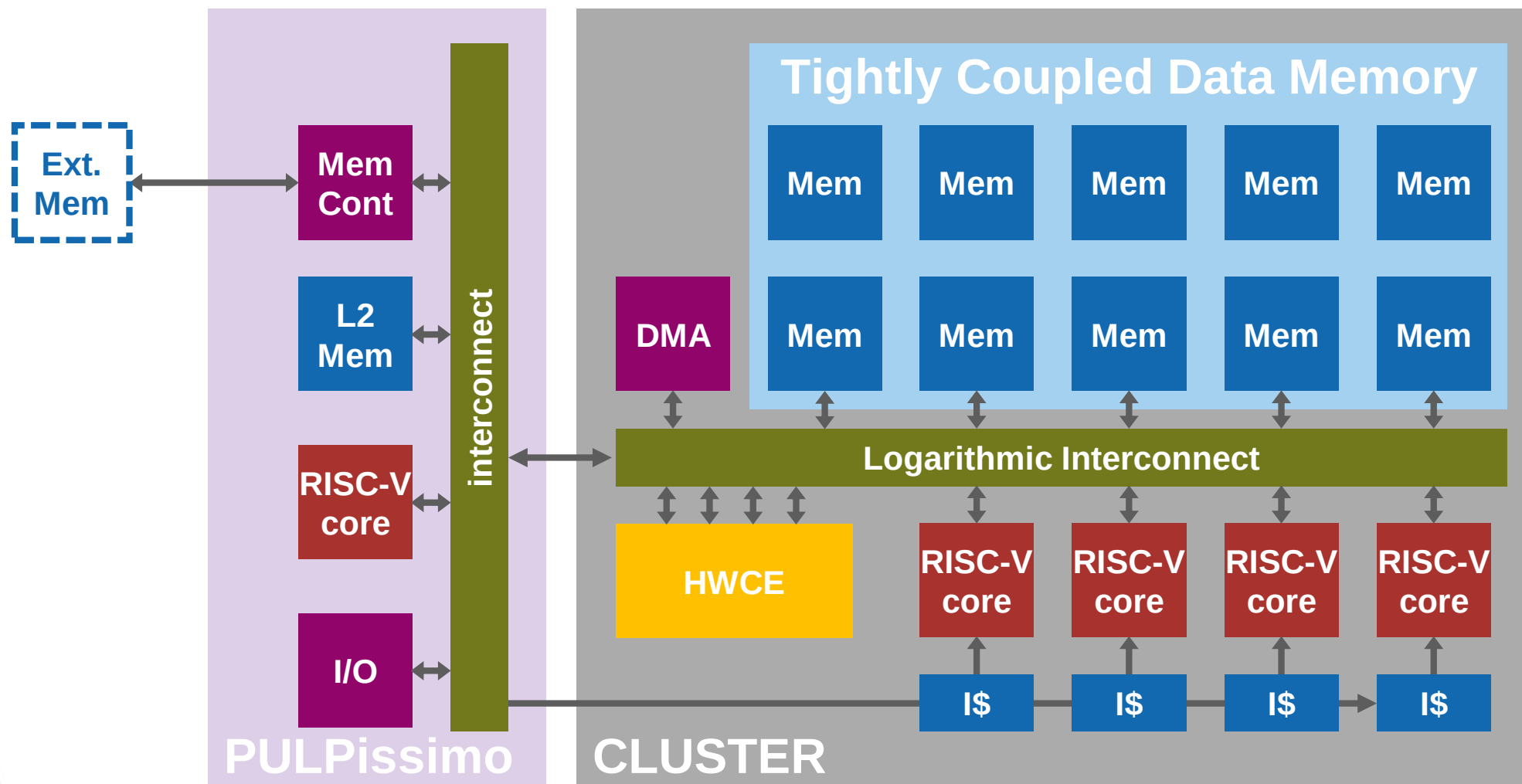
asimovinstitute.org/neural-network-zoo



Hardware Processing Engines (HWPEs)



Sub-pJ/W Accelerator; Tightly-coupled HW Compute Engine



Acceleration with flexibility: zero-copy HW-SW cooperation

More Efficiency (2): Extreme Quantization

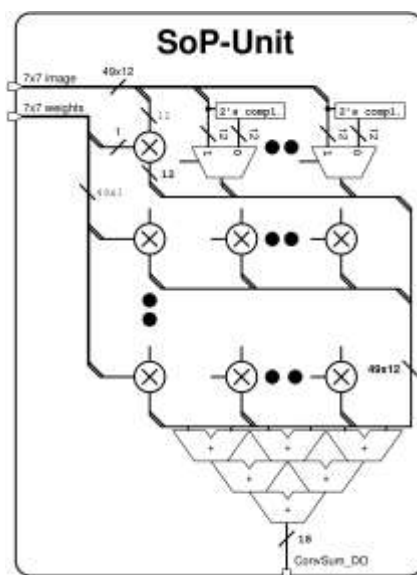
Low(er) precision: 8→4→2

Model	Bit-width	Top-1 error
ResNet-18 ref	32	31.73%
INQ	5	31.02%
INQ	4	31.11%
INQ	3	31.92%
INQ	2 (ternary)	33.98%

SOA INQ retraining

2.2% loss → 0% with 20% larger net

MULT → MUX



Equivalent for 7x7 SoP

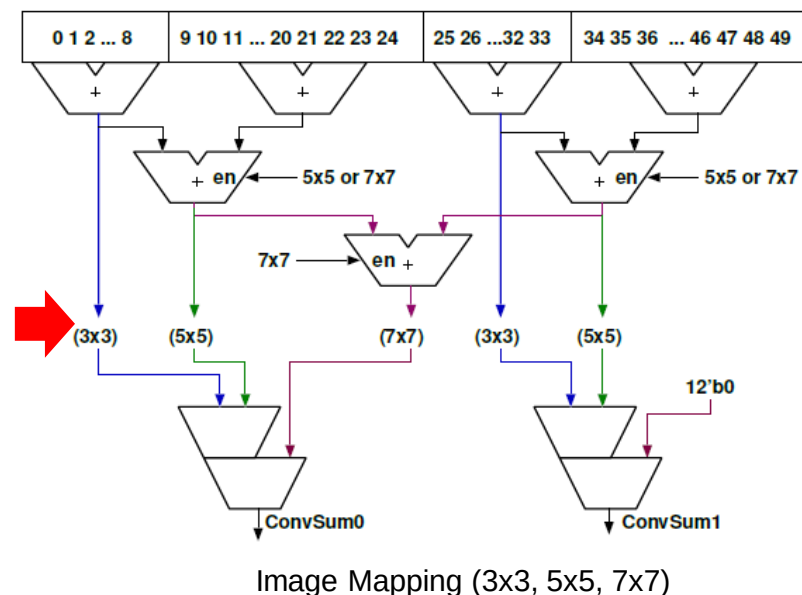


Image Mapping (3x3, 5x5, 7x7)

1 MAC Op = 2 Op (1 Op for the “sign-reverse”, 1 Op for the add).

From +/-1 Binarization to XNORs

$$y(k_{out}) = \text{binarize}_{\pm 1} \left(\mathbf{b}_{k_{out}} + \sum_{k_{in}} \left(\mathbf{W}(k_{out}, k_{in}) \otimes \mathbf{x}(k_{in}) \right) \right)$$

XNOR

$$\text{binarize}_{\pm 1}(t) = \text{sign} \left(\gamma \frac{t - \mu}{\sigma} + \beta \right)$$

$$\text{binarize}_{0,1}(t) = \begin{cases} 1 & \text{if } t \geq -\kappa/\lambda \doteq \tau, \text{ else } 0 & (\text{when } \lambda > 0) \\ 1 & \text{if } t \leq -\kappa/\lambda \doteq \tau, \text{ else } 0 & (\text{when } \lambda < 0) \end{cases}$$

$$y(k_{out}) = \text{binarize}_{0,1} \left(\sum_{k_{in}} \left(\mathbf{W}(k_{out}, k_{in}) \otimes \mathbf{x}(k_{in}) \right) \right)$$

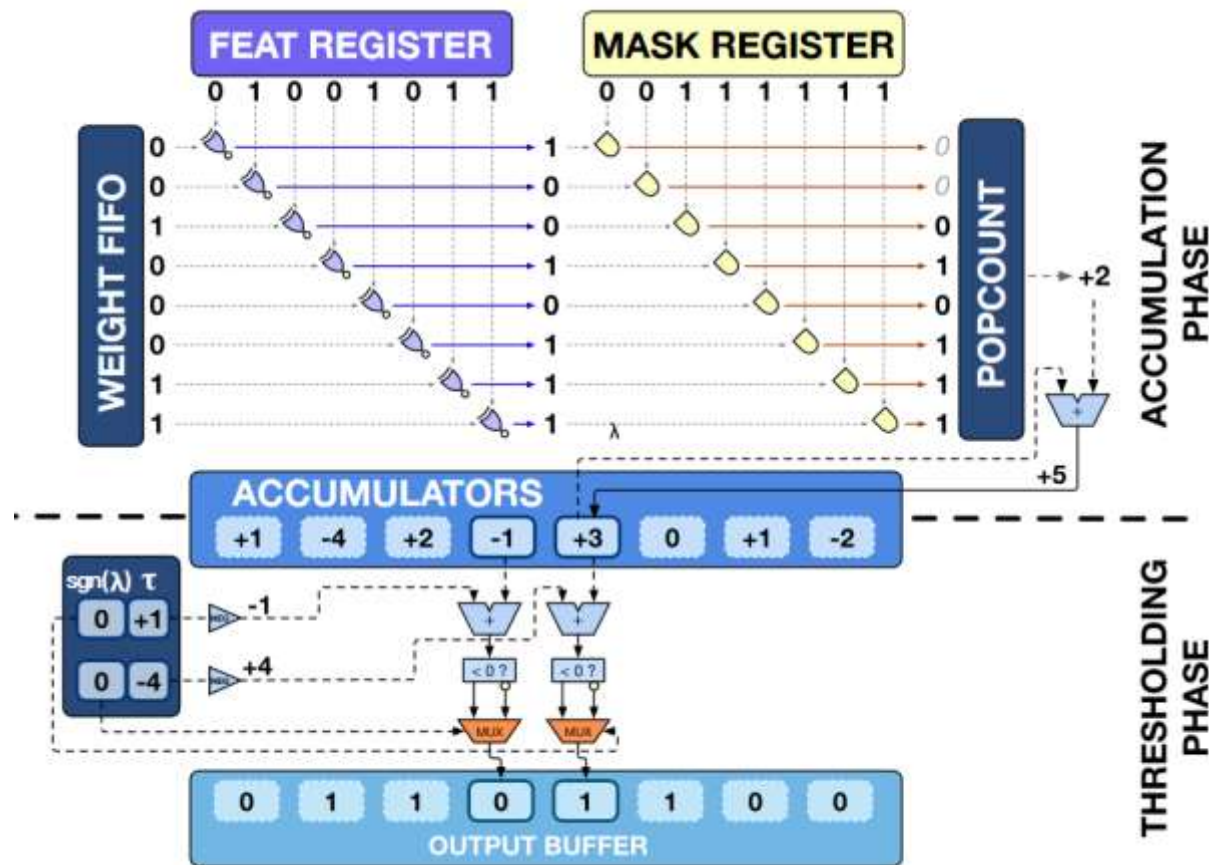
Thresholding

Multi-bit accumulation

Binary product → XOR

A	B	out	A	B	out
-1	-1	+1	0	0	1
-1	+1	-1	0	1	0
+1	-1	-1	1	0	0
+1	+1	+1	1	1	1

XNE: XNOR Neural Engine



BINCONV: Binary dot-product and thresholding logic array

XNE Energy Efficiency

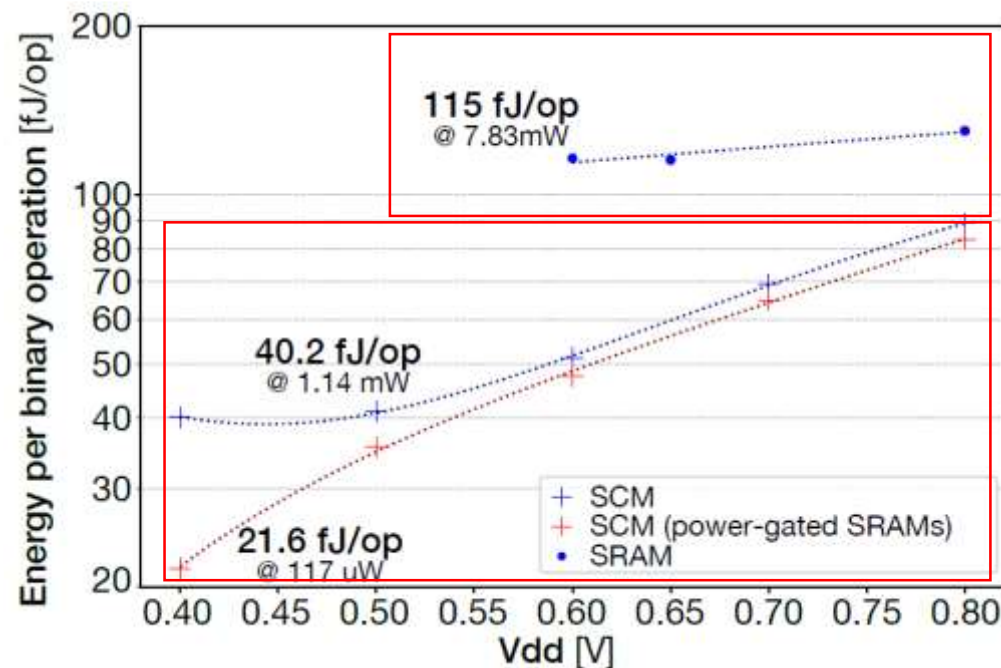
22 FDX measured silicon

With SRAMs, max eff
@ 0.65V **8.7 Top/s/W**

With SCMs, max eff
@ 0.5V **46.3 Top/s/W**



100+ TOPS/W (10fJ/OP)
Is reachable!



But... Accuracy Loss is high even with retraining (10%+)
Need flexible precision tuning!

Flexibility needed: Binary-Based Quantization (BBQ)

QNN layer :

$$y(k_{out}) = \text{quant} \left(\sum_{k_{in}} \overbrace{(\mathbf{W}(k_{out}, k_{in}) \otimes \mathbf{x}(k_{in}))}^{\text{INT32 accumulator}} \right)$$

Q-bit output fmaps

M-bit weights

N-bit input fmaps

Many $M \times N$ bits products...

... but one $M \times N$ product is the superposition of $M \times N$ **1-bit** products!

$$y(k_{out}) = \text{quant} \left(\sum_{i=0..M} \sum_{j=0..N} \sum_{k_{in}} \overbrace{2^i 2^j}^{\text{power-of-2 scaling factors}} (\underbrace{\mathbf{w}_{bin}(k_{out}, k_{in})}_{\text{1-bit weights}} \otimes \underbrace{\mathbf{x}_{bin}(k_{in})}_{\text{1-bit input fmaps}}) \right)$$

Q-bit output fmaps

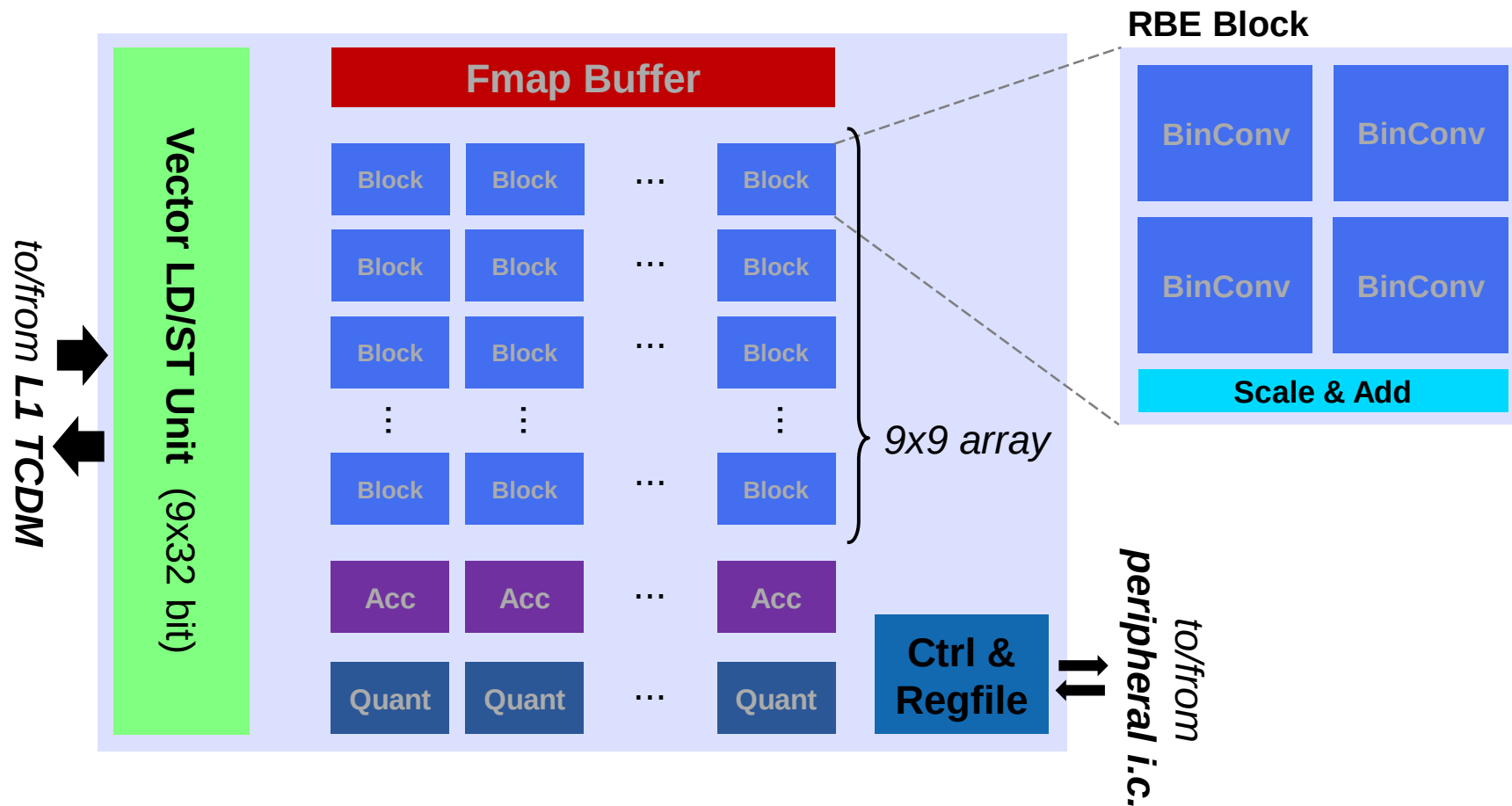
1-bit weights

1-bit input fmaps

One quantized NN can be emulated by superposition of power-of-2 weighted $M \times N$ binary NN

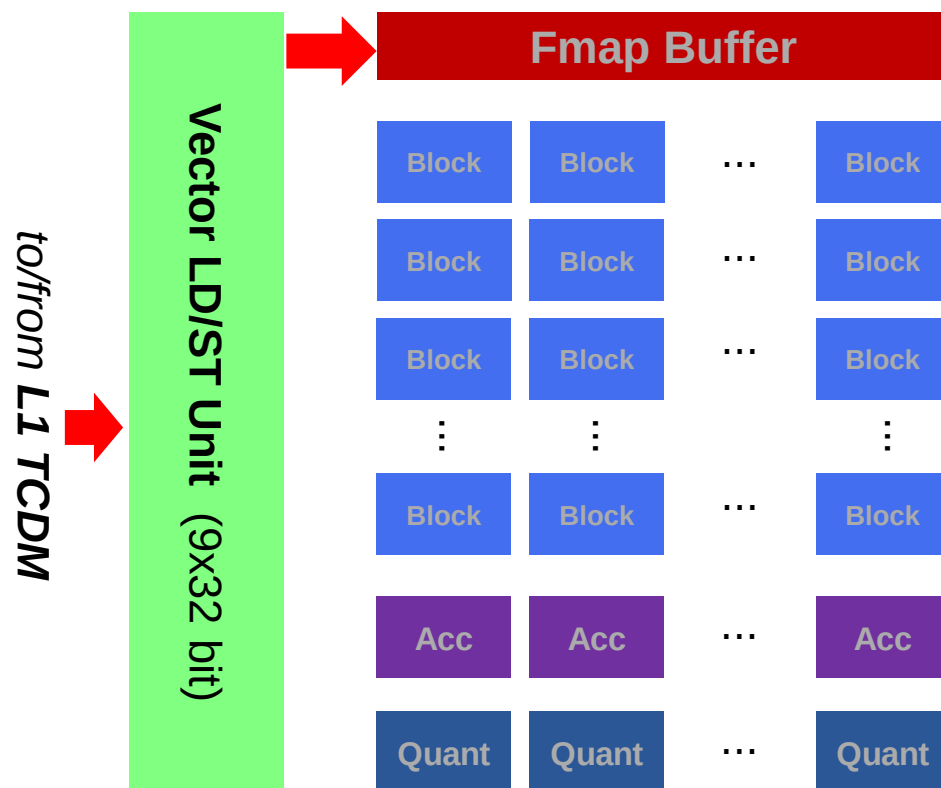
Reconfigurable Binary Engine

$$y(k_{out}) = \text{quant} \left(\sum_{i=0..M} \sum_{j=0..N} \sum_{k_{in}} 2^i 2^j (\mathbf{W}_{\text{bin}}(k_{out}, k_{in}) \otimes \mathbf{x}_{\text{bin}}(k_{in})) \right) \quad \text{e.g. a 3x3 conv with } \mathbf{N}=4 \text{ bits}$$



Reconfigurable Binary Engine

$$y(k_{out}) = \text{quant} \left(\sum_{i=0..M} \sum_{j=0..N} \sum_{k_{in}} 2^i 2^j (\mathbf{w}_{bin}(k_{out}, k_{in}) \otimes \mathbf{x}_{bin}(k_{in})) \right)$$



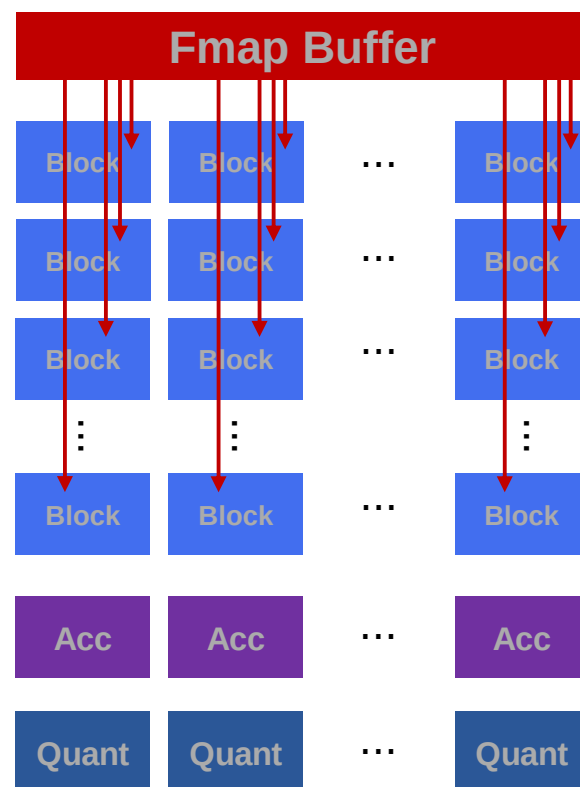
1. preload **fmap buffer** with input activations (up to 5x5 pixels, 32 channels, **N**=4 bits/pixel)

Reconfigurable Binary Engine

$$y(k_{out}) = \text{quant} \left(\sum_{i=0..M} \sum_{j=0..N} \sum_{k_{in}} 2^i 2^j (\mathbf{w}_{bin}(k_{out}, k_{in}) \otimes \mathbf{x}_{bin}(k_{in})) \right)$$

to/from L1 TCDM

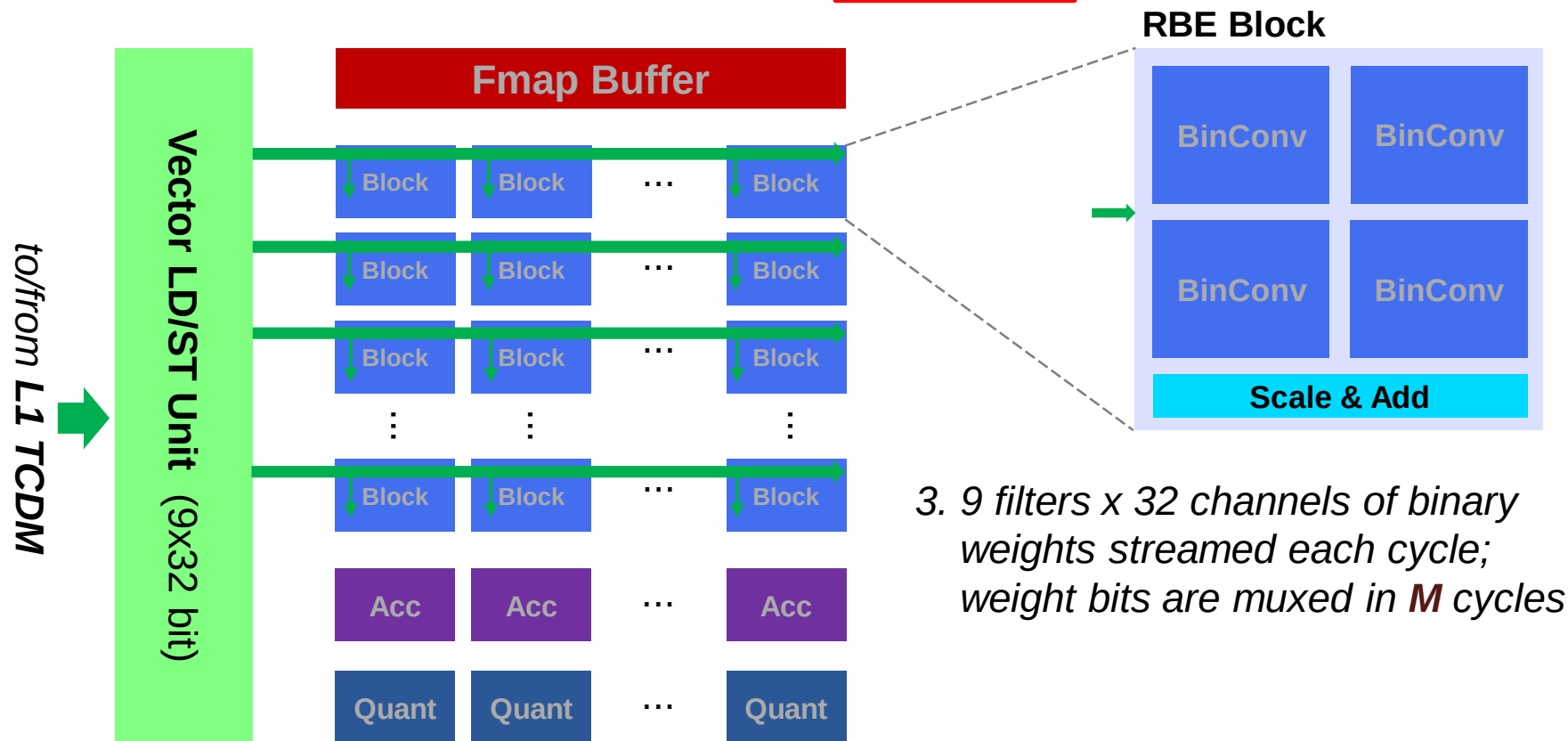
Vector LD/ST Unit (9x32 bit)



1. preload **fmap buffer** with input activations (up to 5x5 pixels, 32 channels, 4 bits/pixel)
2. each of the 81 blocks has as input one of the 25 px (32x4 bits)

Reconfigurable Binary Engine

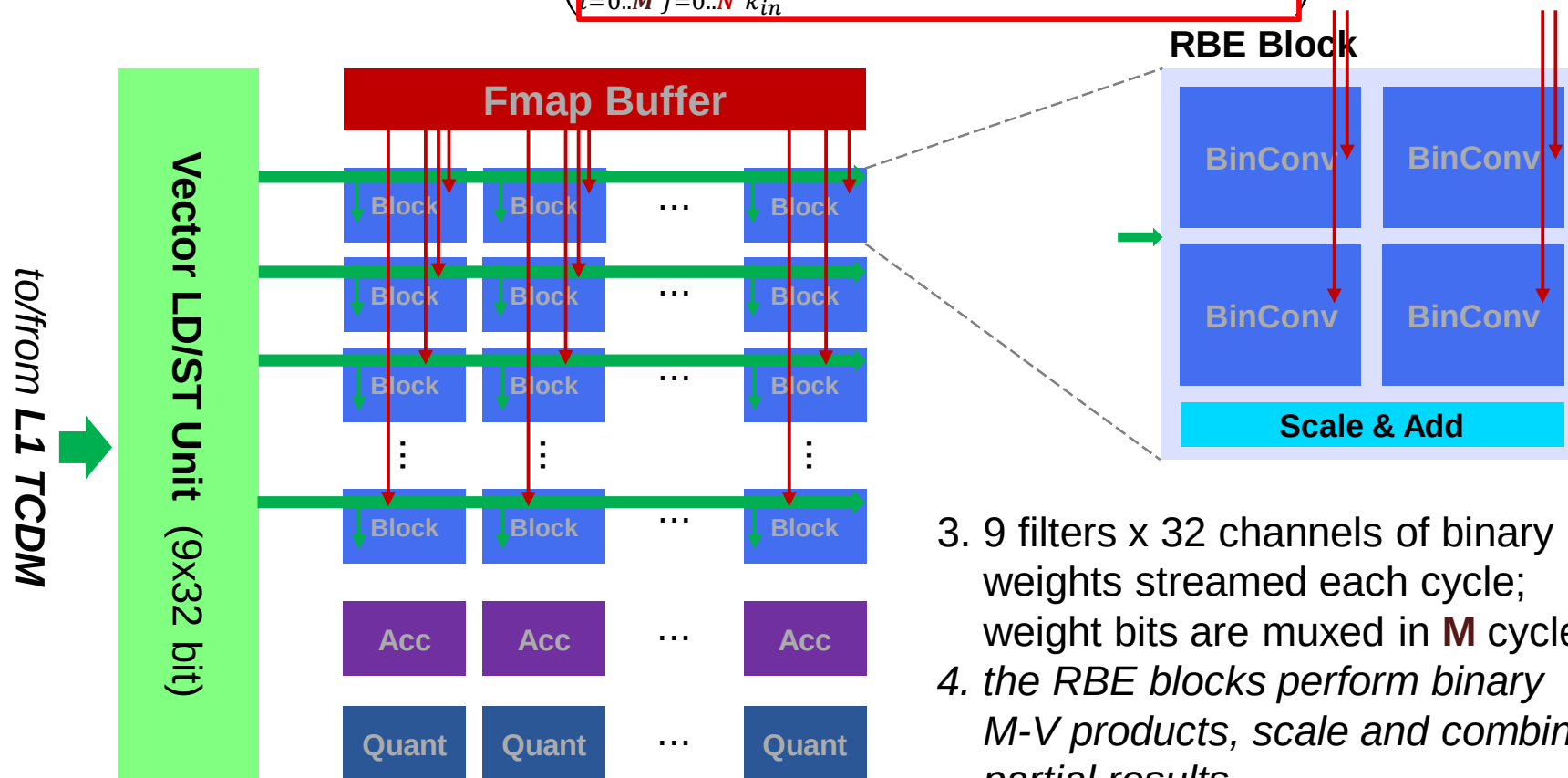
$$y(k_{out}) = \text{quant} \left(\sum_{i=0..M} \sum_{j=0..N} \sum_{k_{in}} 2^i 2^j (\mathbf{W}_{\text{bin}}(k_{out}, k_{in}) \otimes \mathbf{x}_{\text{bin}}(k_{in})) \right)$$



3. 9 filters x 32 channels of binary weights streamed each cycle; weight bits are muxed in **M** cycles

Reconfigurable Binary Engine

$$y(k_{out}) = \text{quant} \left(\sum_{i=0..M} \sum_{j=0..N} \sum_{k_{in}} 2^i 2^j (\mathbf{W}_{bin}(k_{out}, k_{in}) \otimes \mathbf{x}_{bin}(k_{in})) \right)$$



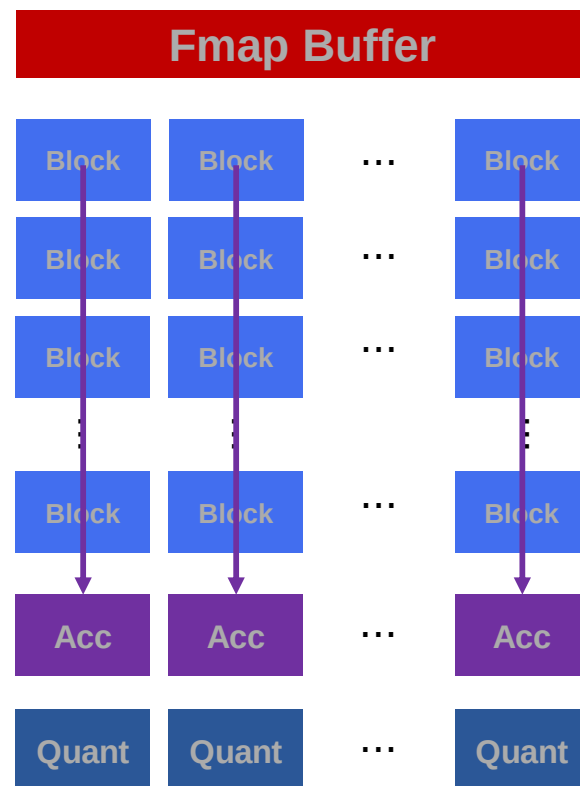
3. 9 filters x 32 channels of binary weights streamed each cycle; weight bits are muxed in **M** cycles
4. the RBE blocks perform binary *M-V* products, scale and combine partial results

Reconfigurable Binary Engine

$$y(k_{out}) = \text{quant} \left(\sum_{i=0..M} \sum_{j=0..N} \sum_{k_{in}} 2^i 2^j (\mathbf{w}_{bin}(k_{out}, k_{in}) \otimes \mathbf{x}_{bin}(k_{in})) \right)$$

to/from L1 TCDM

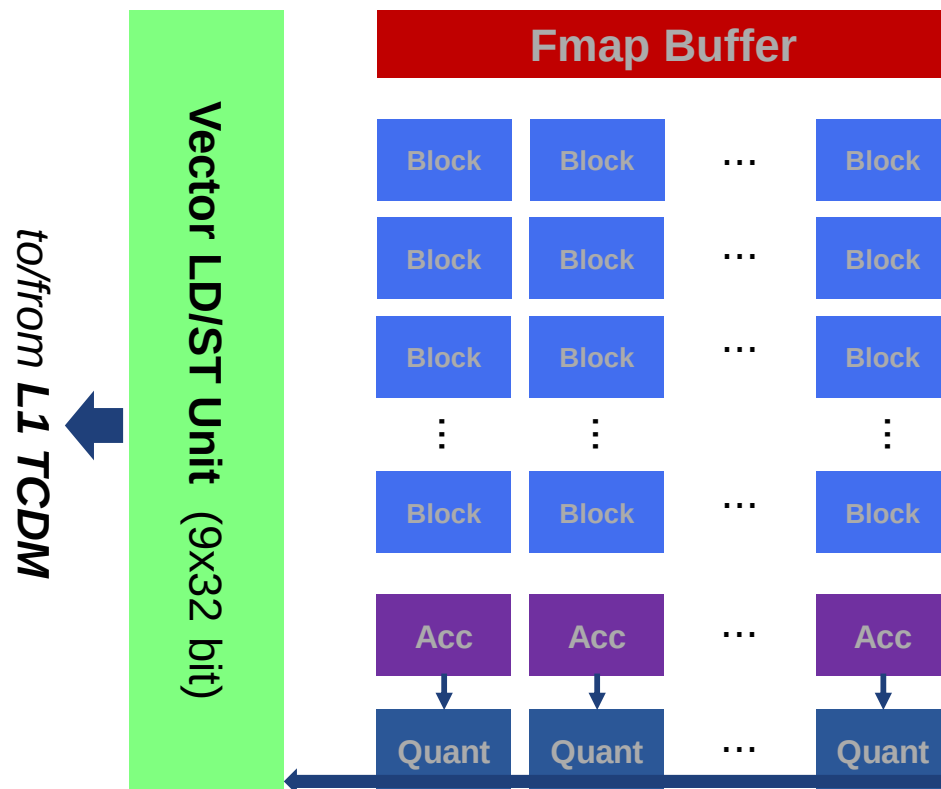
Vector LD/ST Unit (9x32 bit)



5. results are added column-wise and accumulated in 32-bit *accumulator* banks (9x32x32 bits) for many iterations until full accumulation

Reconfigurable Binary Engine

$$y(k_{out}) = \text{quant} \left(\sum_{i=0..M} \sum_{j=0..N} \sum_{k_{in}} 2^i 2^j (\mathbf{w}_{bin}(k_{out}, k_{in}) \otimes \mathbf{x}_{bin}(k_{in})) \right)$$

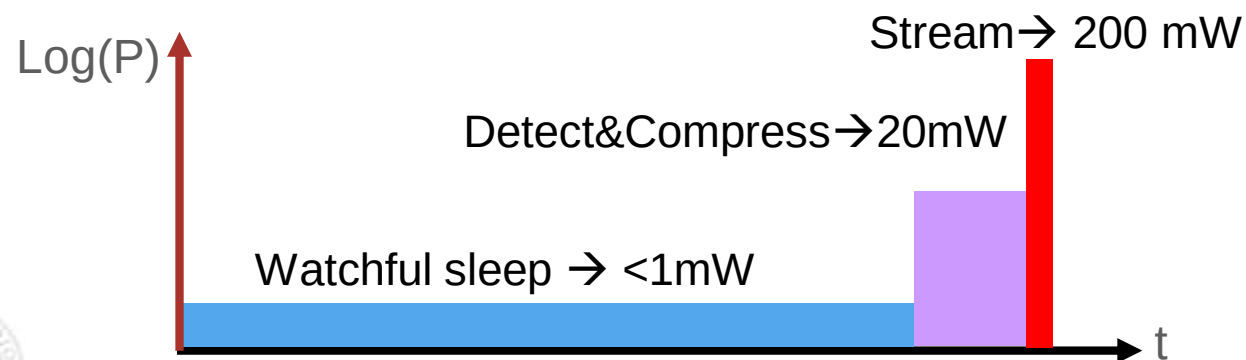
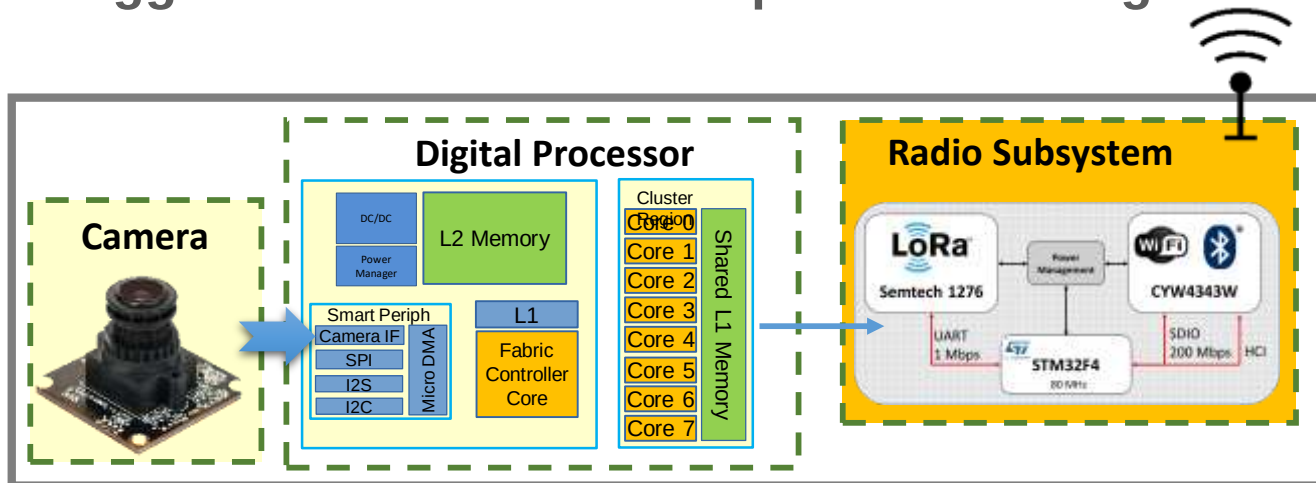


5. results are added column-wise and accumulated in 32-bit **accumulator** banks (9x32x32 bits) for many iterations until full accumulation
6. after full accumulation, the accumulator values are **quantized** and **streamed out**

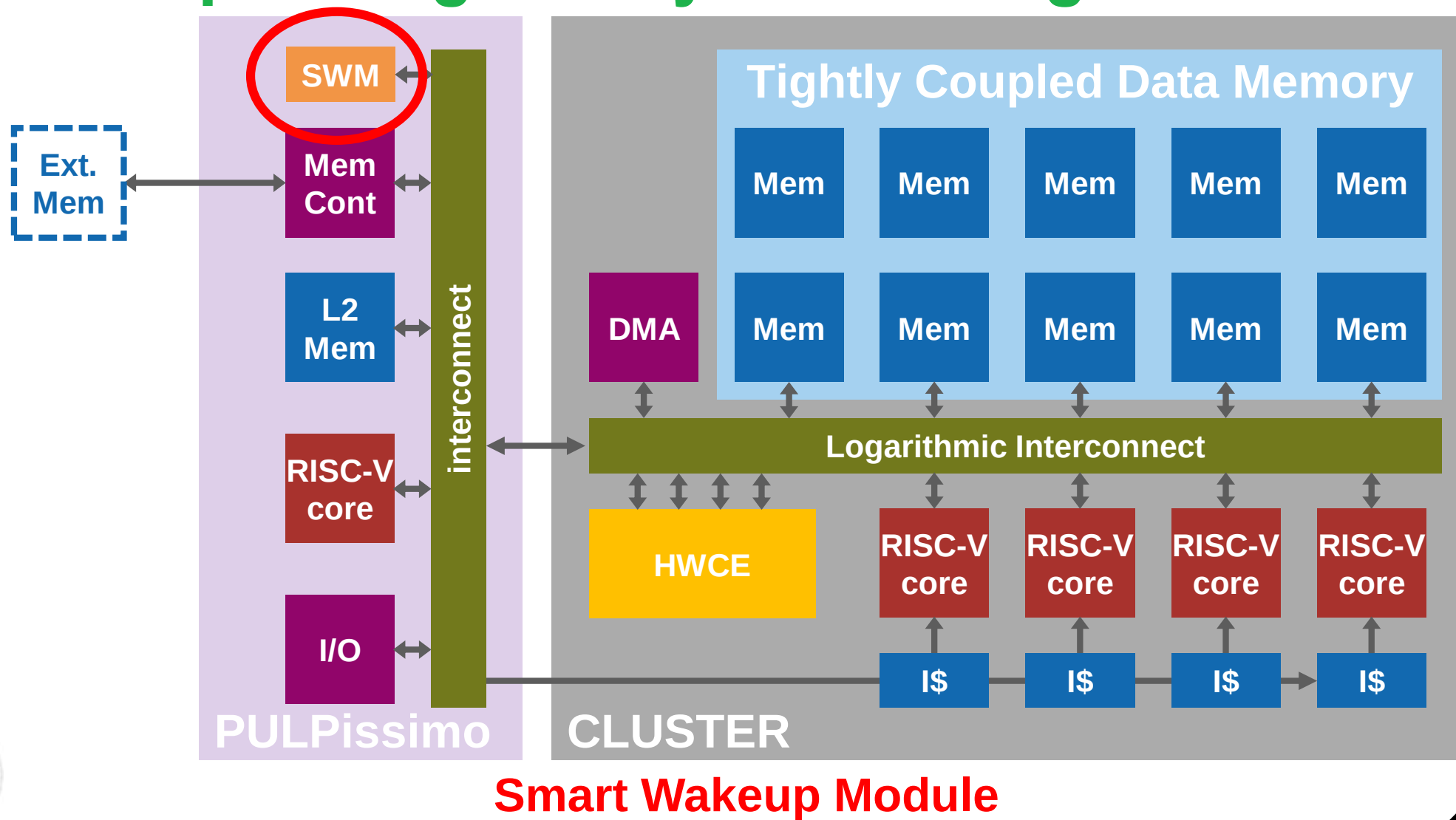
What about μW «sleep»?

Small always-on network $\rightarrow$

Triggers alarm and video capture/streaming for cloud-based forensics

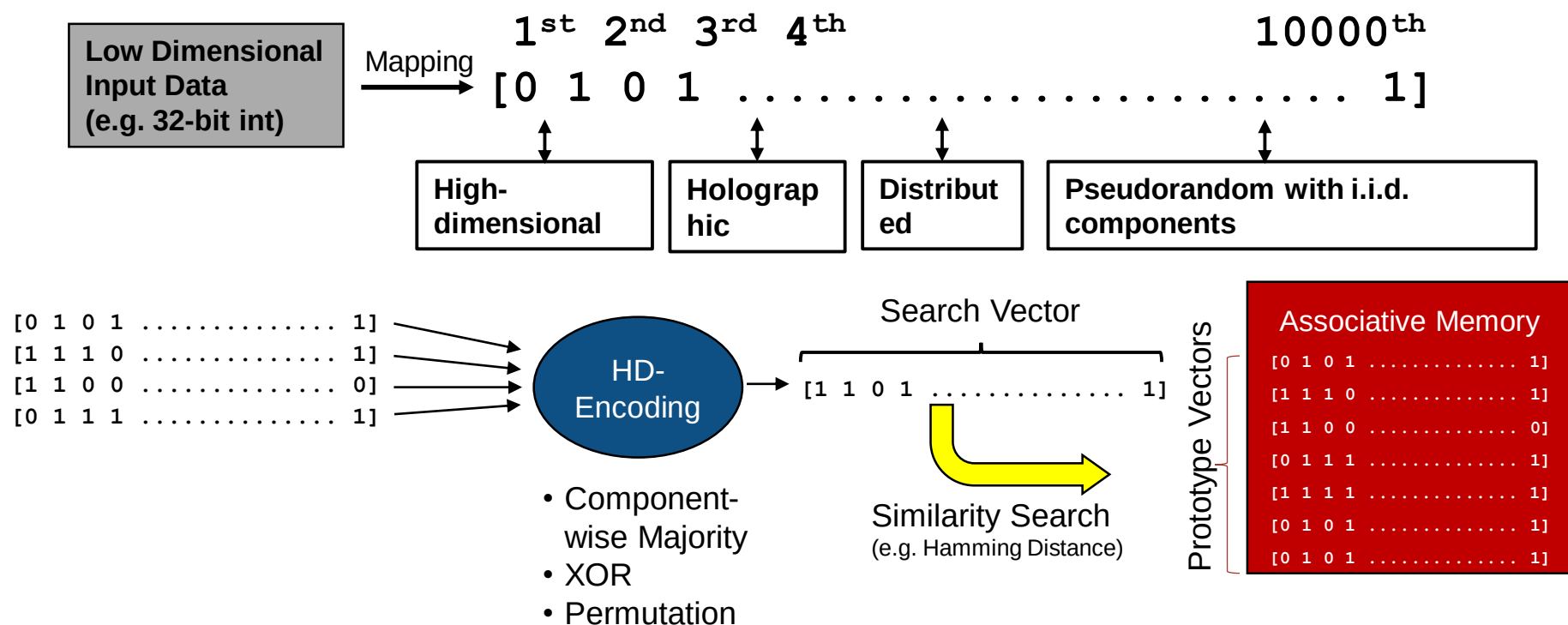


Need μ W-range always-on Intelligence



Smart Wakeup Module

Not Only CNNs: Hyper-Dimensional Computing



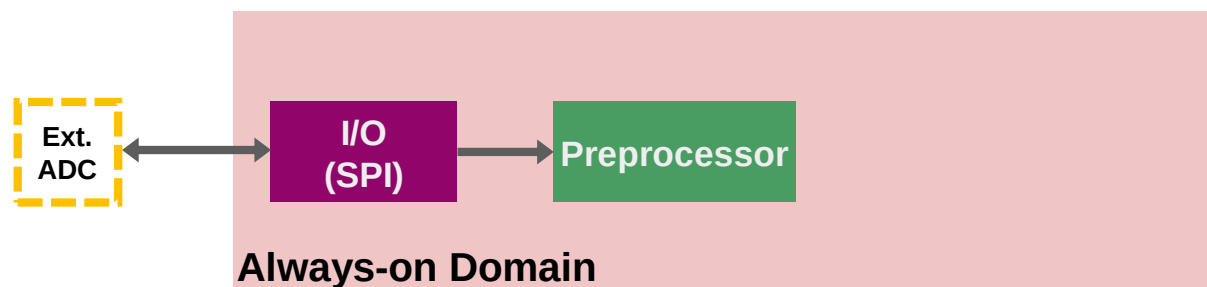
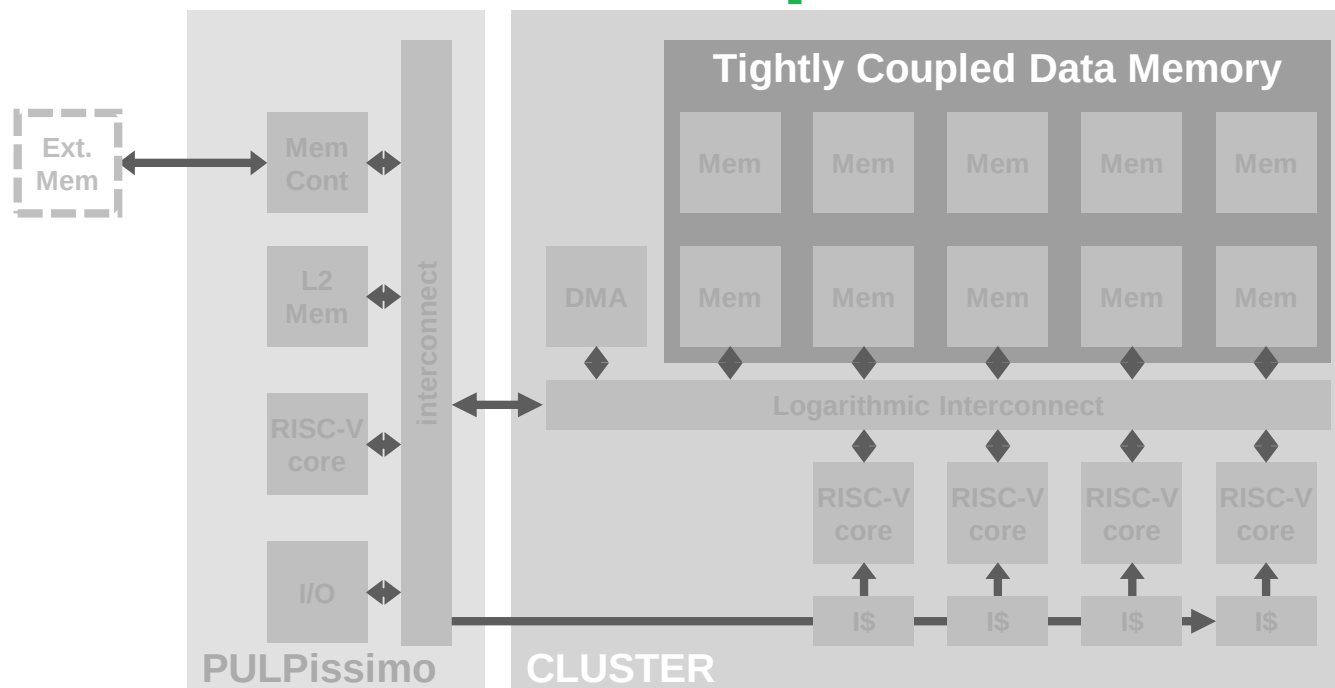
Highly parallel, fault-tolerant binary operators, assoc-min-distance search



Merge storage & computation
i.e. **In-memory computing**

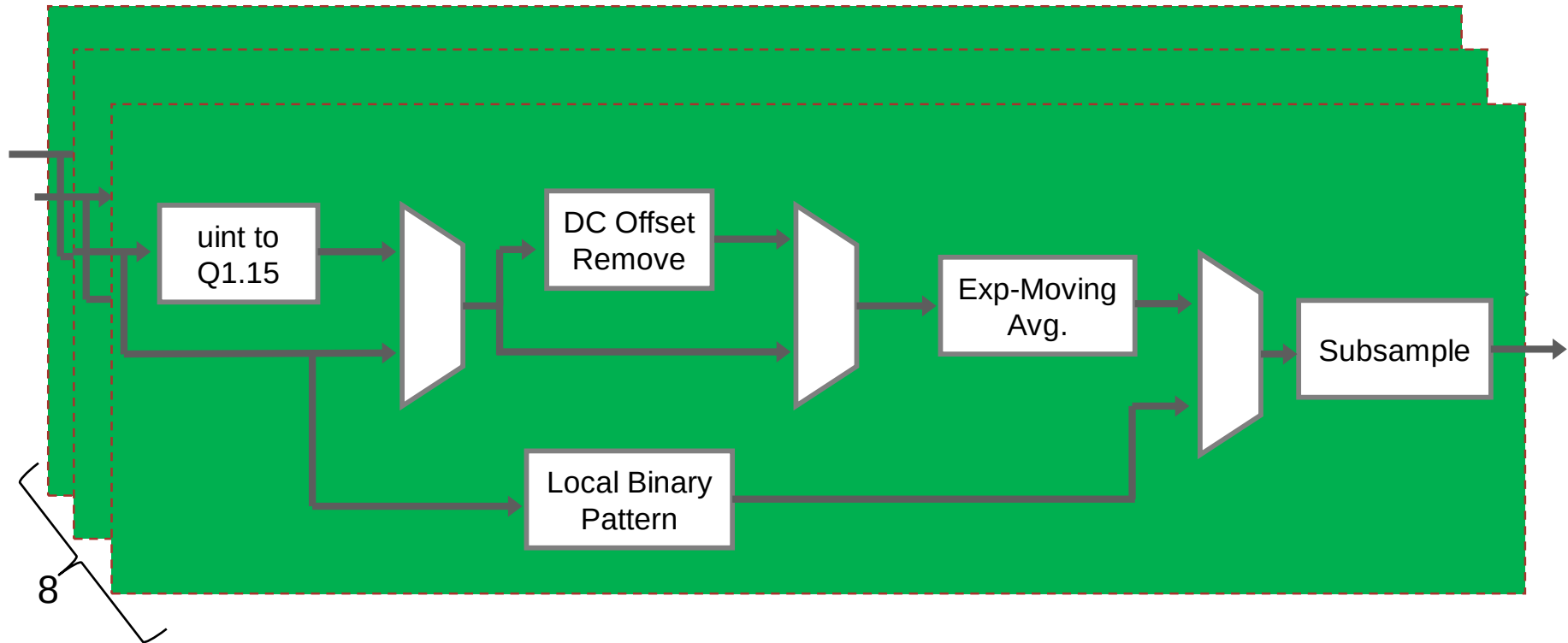


HD-Based smart Wake-Up Module



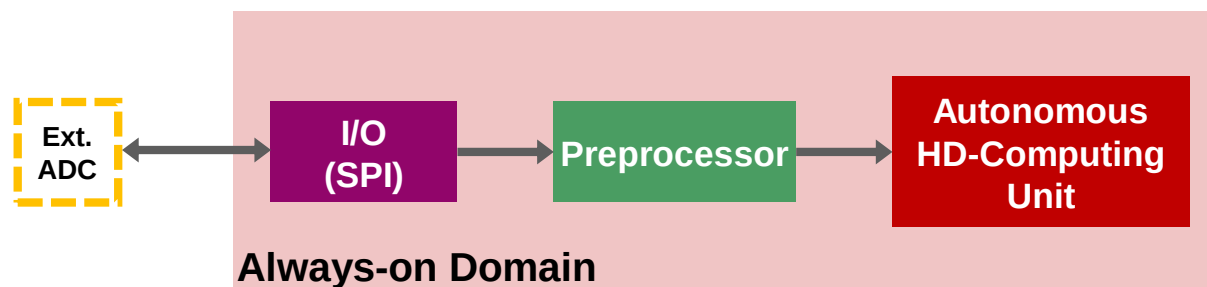
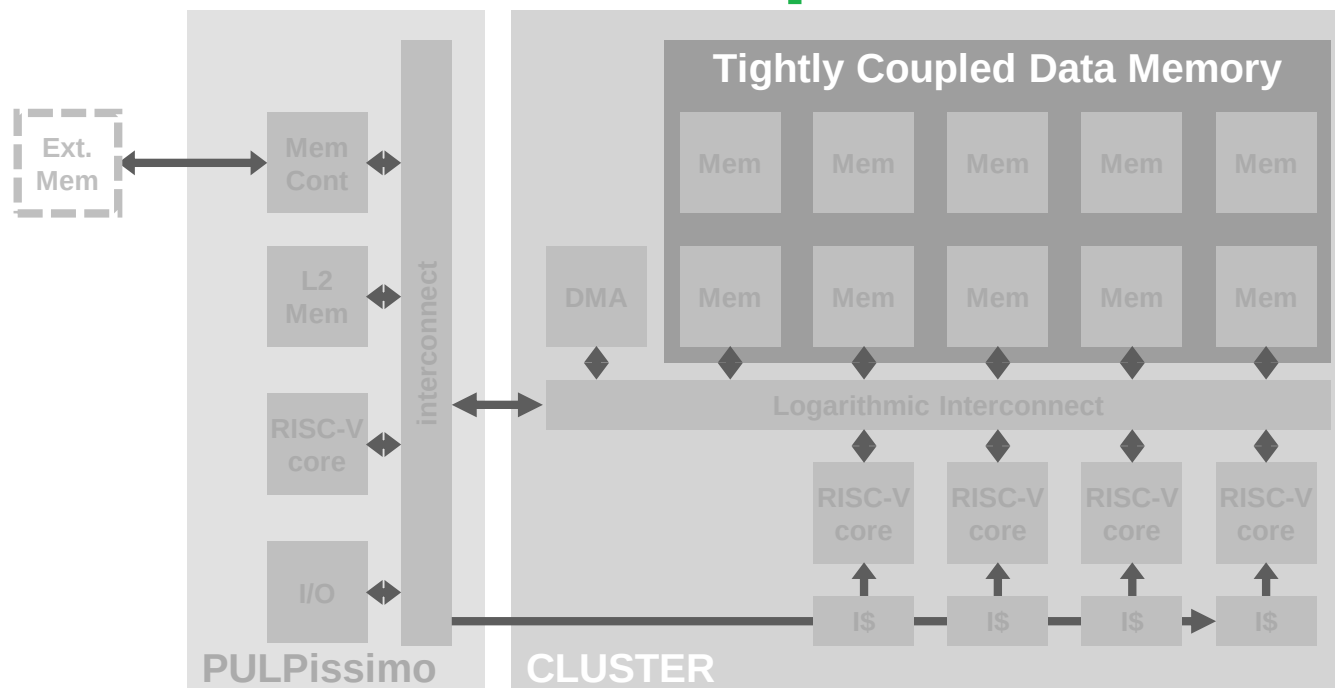


Preprocessor



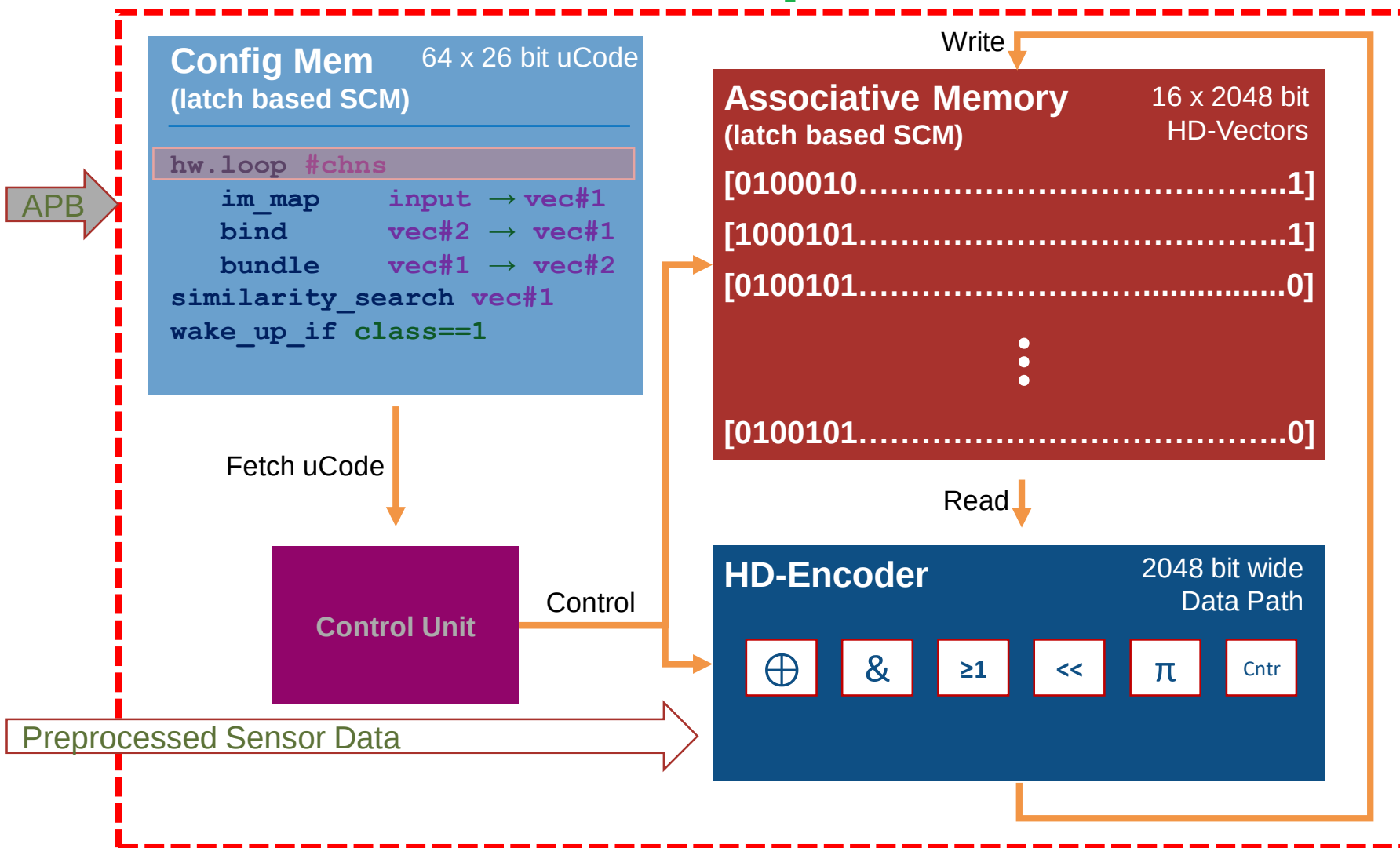


HD-Based smart Wake-Up Module

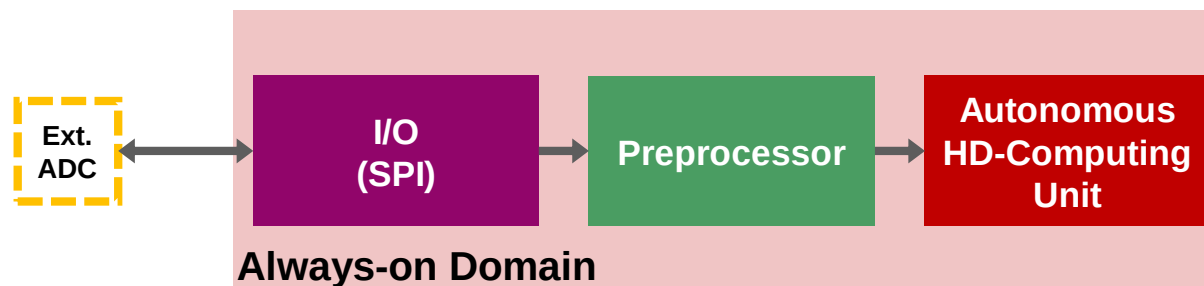
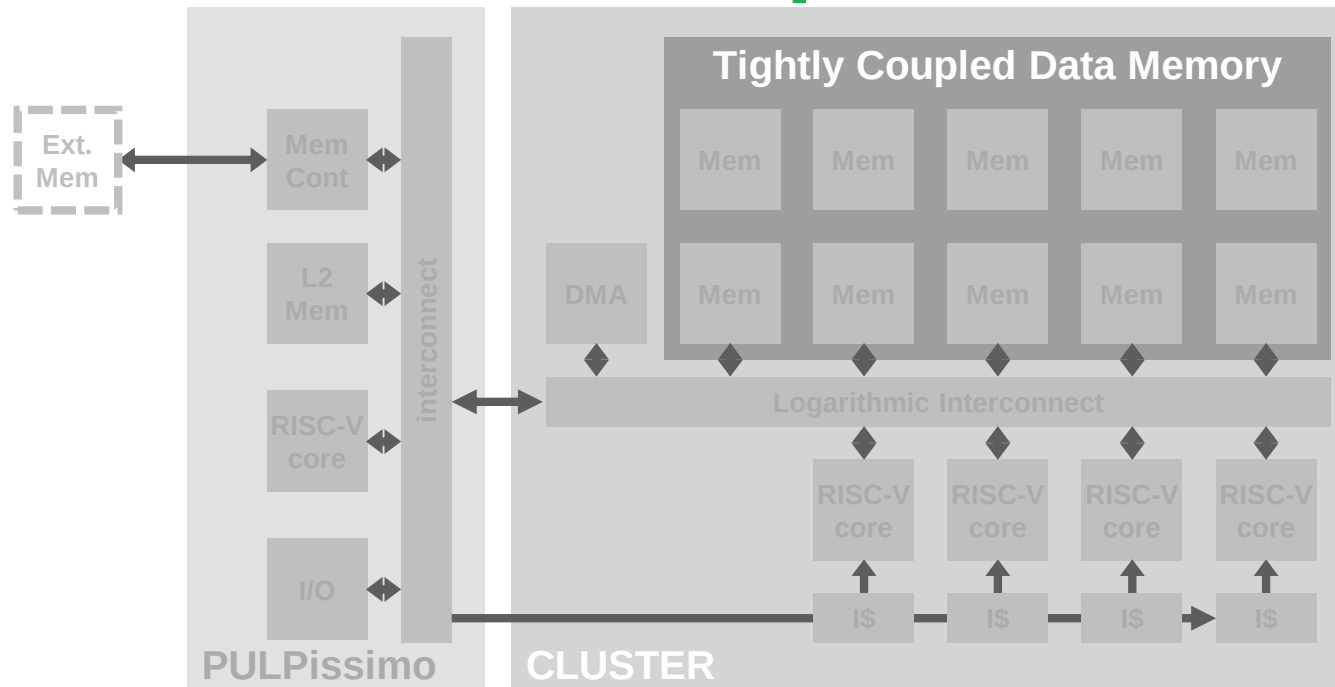




HD-Based smart Wake-Up Module

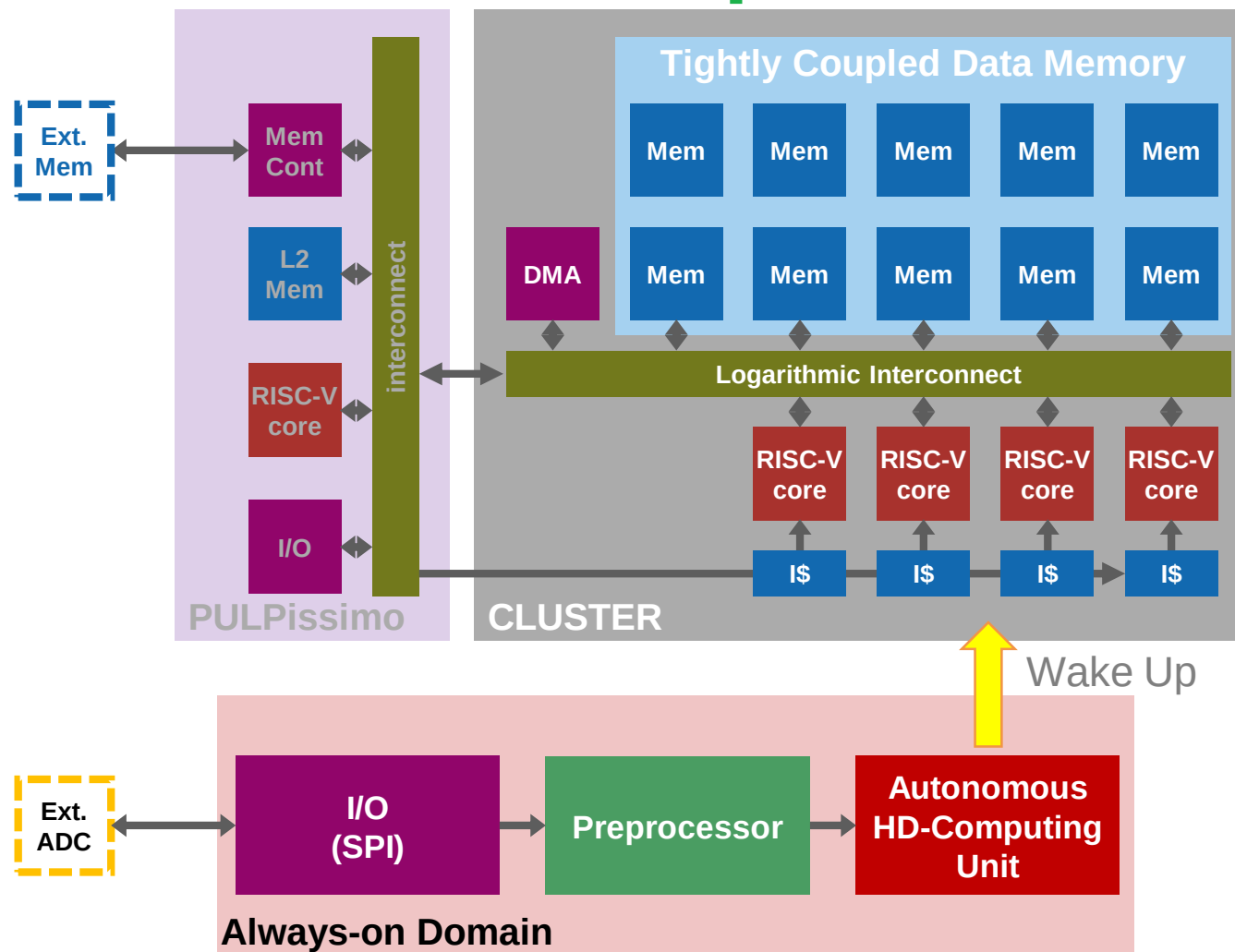


HD-Based smart Wake-Up Module





HD-Based smart Wake-Up Module





HD-Based smart Wake-Up Module

Results (post-P&R) - TO done

Technology	GF22 UHT
Area	670kGE
Max. Frequency	3 MHz
SCM-Memory	32 kBit
Module Power Consumption (@ 1 kSPS/channel, 3 channels)	~ 15uW

Integrating Accelerators into PULP – The big picture

RISC-V Cores

RI5CY

32b

Ibex

32b

Snitch

32b

Ariane
+ Ara
64b

Peripherals

JTAG

UART

DMA

SPI

I2S

GPIO

Interconnect

Logarithmic interconnect

APB – Peripheral Bus

AXI4 – Interconnect

Platforms

Tens of active users, many use-cases
HW, SW specialization, verification, documentation, training
Cannot be sustained by one University, or two...

Single Core

- PULPino
- PULPissimo

RV

Multi-core

- Fulmine
- Mr. Wolf

R5

Multi-cluster

- Hero
- Open Piton

IOT

HPC

Accelerators

HWCE
(convolution)

Neurostream
(ML)

HWCrypt
(crypto)

PULPO
(1st ord. opt)

Academic open source → Industrial open source



OPENHW GROUP™
— PROVEN PROCESSOR IP —

Rick O'Connor (OpenHW CEO, former RISC-V foundation director)

- **OpenHW Group** is a not-for-profit, global organization (EU,NA,Asia) driven by its members and individual contributors where HW and SW designers collaborate in the development of open-source cores, related IP, tools and SW such as the **Core-V** family of cores.
- **OpenHW Group** provides an infrastructure for hosting high quality open-source HW developments in line with industry best practices.



RI5CY, ARIANE



A vertical, application-focused open approach



- OpenTitan is **the first open source** silicon project building a transparent, high-quality reference design for silicon root of trust (RoT) chips.
- Founding Partners



Giesecke+Devrient
Creating Confidence



Feel the momentum!

Ibex RISC-V core, flash interface, communications ports, cryptography accelerators, and more.

Vibrant repository

35+

Contributors

1300

+

Contributions

470

GitHub Issues



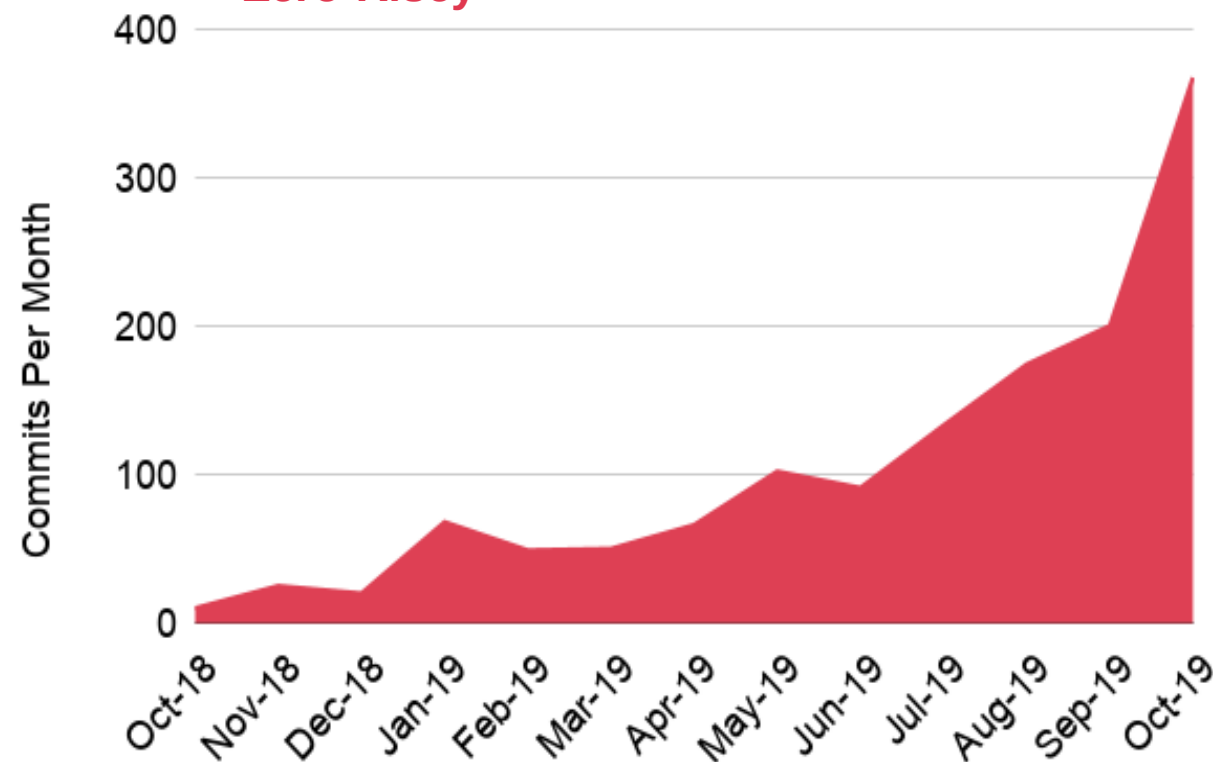
PULP
Parallel Ultra Low Power



lowRISC

Zero-Riscy

Ibex



HPC Vertical: The European Processor Initiative



Europe Needs its own Processors

- Processors now control almost every aspect of our lives
- Security (back doors etc.)
- Possible future restrictions on exports to EU due to increasing protectionism
- A competitive EU supply chain for HPC technologies will create jobs and growth in Europe
- Sovereignty (data, economical, embargo)



- High Performance General Purpose Processor for HPC
- High-performance RISC-V based accelerator**
- Computing platform for autonomous cars
- Will also target the AI, Big Data and other markets in order to be economically sustainable



Closing thoughts... the open HW revolution

For science ... fundamental “research infrastructure”

- Community building: sharing of ideas, artefacts
- Fantastic tool for dissemination (more citations 😊)!
- Reduce “getting up to speed” overhead for partners
- Enables fair and well controlled benchmarking

For Business ... it is truly disruptive

- Reduces the NRE cost for silicon design
- Faster innovation path for startups
- New business models (for profit and non-for profit)
- Helps exchange of information across NDA walls
- Great for Marketing & Training

For society ...long term sustained benefits

- More innovation, growth, jobs
- Personalized silicon vision “Moore-for-all”
- More Secure, safe, auditable HW



Posh Open Source Hardware (**POSH**):

An open source System on Chip (SoC) design and verification ecosystem that enables cost effective design of ultra-complex SoCs

USA's
electronics resurgence initiative



PULP

Parallel Ultra Low Power

Luca Benini, Davide Rossi, Andrea Borghesi, Michele Magno, Simone Benatti, Francesco Conti, Francesco Beneventi, Daniele Palossi, Giuseppe Tagliavini, Antonio Pullini, Germain Haugou, Lukas Cavigelli, Manuele Rusci, Florian Glaser, Renzo Andri, Fabio Montagna, Bjoern Forsberg, Pasquale Davide Schiavone, Alfio Di Mauro, Victor Javier Kartsch Morinigo, Tommaso Polonelli, Fabian Schuiki, Stefan Mach, Andreas Kurth, Florian Zaruba, Manuel Eggimann, Philipp Mayer, Marco Guermandi, Xiaying Wang, Michael Hersche, Robert Balas, Antonio Mastrandrea, Matheus Cavalcante, Angelo Garofalo, Alessio Burrello, Gianna Paulin, Georg Rutishauser, Andrea Cossettini, Luca Bertaccini, Maxim Mattheeuws, Samuel Riedel, Sergei Vostrikov, Vlad Niculescu, Frank K. Gurkaynak, *and many more that we forgot to mention*



<http://pulp-platform.org>



@pulp_platform

