

OctopusScheduler: On-Chip DNN Scheduling on the SpiNNaker2 Neuromorphic MPSoC

Tim Langer*, Matthias Jobst*[†], Chen Liu*, Florian Kelber*, Bernhard Vogginger*, Christian Mayr*^{†‡}

*Chair of Highly-Parallel VLSI Systems and Neuro-Microelectronics, TU Dresden, Germany

[†]Centre for Tactile Internet with Human-in-the-Loop (CeTI), TU Dresden, Germany [‡]ScaDS.AI Dresden/Leipzig, Germany
{tim.langer, matthias.jobst2, chen.liu, florian.kelber, bernhard.vogginger, christian.mayr}@tu-dresden.de

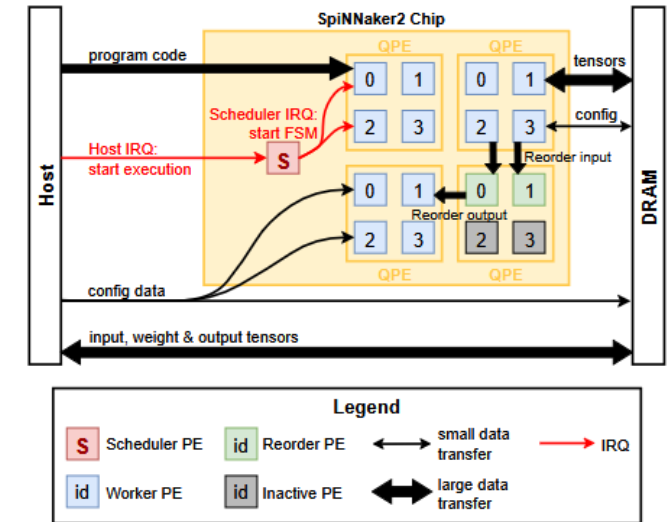
Presenter: Tim Langer

27.03.2025 @ NICE'25 (Heidelberg)

OctopusScheduler Contributions

What have we developed?

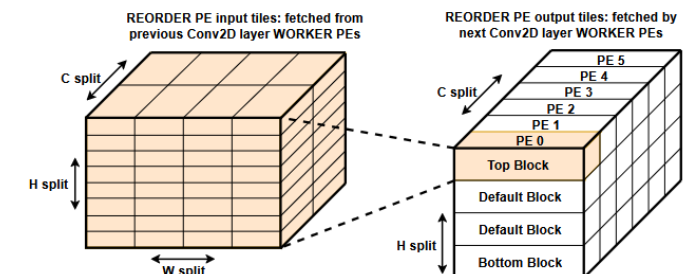
- framework for execution of single DNN layers on SpiNNaker2
- on-chip scheduling for SpiNNaker2:
 - avoiding host-to-chip communication
 - minimizing scheduling overhead
- hardware-aware strategy for automated tiling design space exploration for SpiNNaker2
- on-chip reordering mechanism for Conv2D on SpiNNaker2:
 - avoiding DRAM transfers of intermediate activations



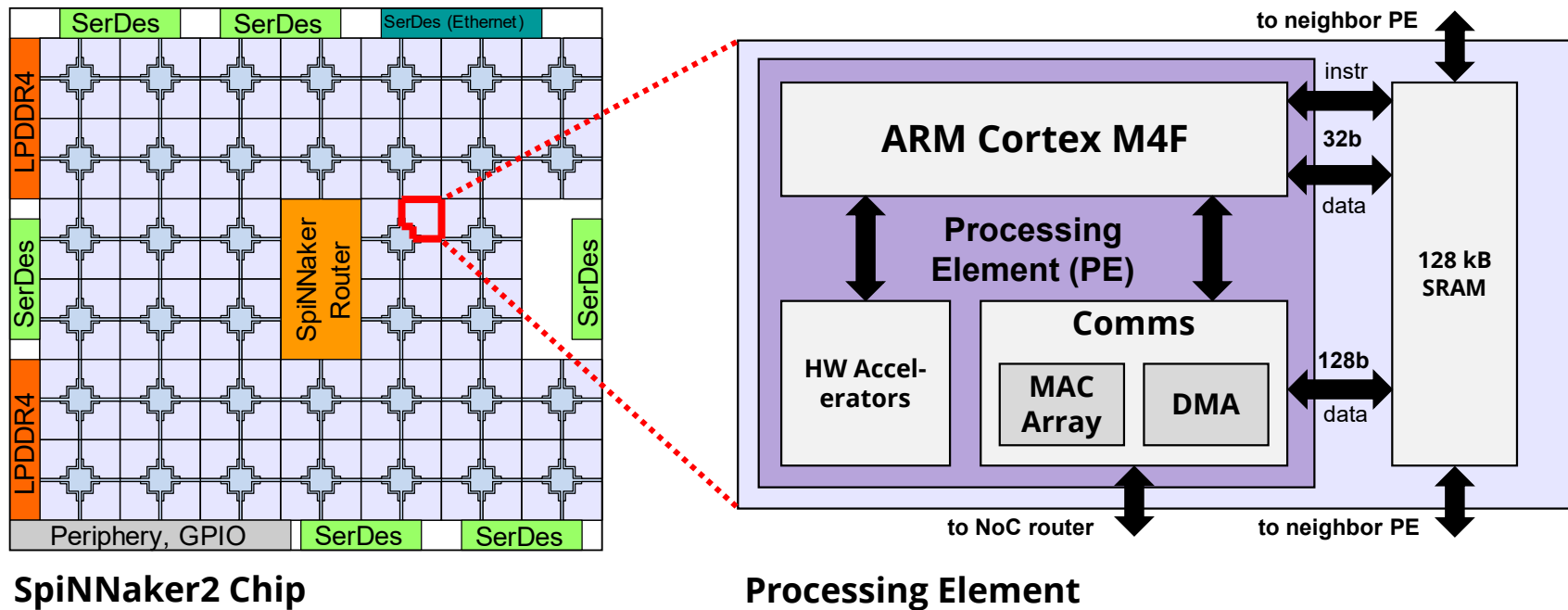
$$\begin{aligned} \mathbf{T}(M) &:= \{m : (m \mid M) \wedge (m \bmod 4 = 0)\} \subset \mathbb{N} \\ \mathbf{T}(P) &:= \{p : (p \mid P) \wedge (p \bmod 16 = 0)\} \subset \mathbb{N} \\ \mathbf{W} &:= \{w : (w \mid (M \cdot P)) \wedge (w \leq \#(\text{PE})_{\text{avail}})\} \subset \mathbb{N} \end{aligned}$$

Tiling Space:

$$\mathbf{S}_{MatMul} = \mathbf{T}(M) \times \mathbf{T}(P) \times \mathbf{W}$$



SpiNNaker2 Hardware Architecture

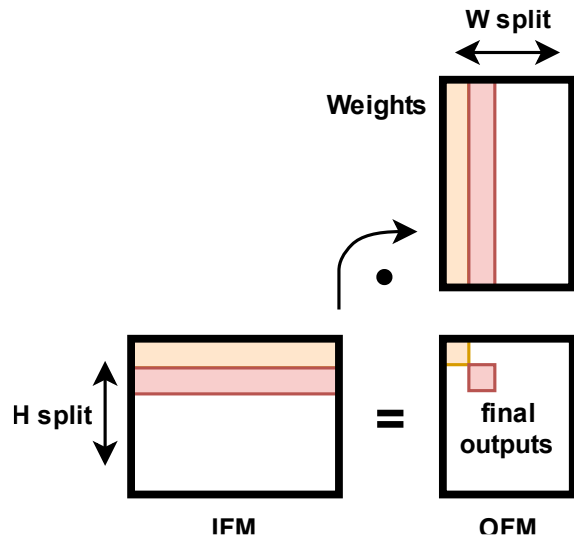


Modified after: Gonzalez, H. A. et al. (2023): „SpiNNaker2: A Large-Scale Neuromorphic System for Event-Based and Asynchronous Machine Learning“. MLNCP Workshop at NeurIPS 2023, New Orleans (USA).

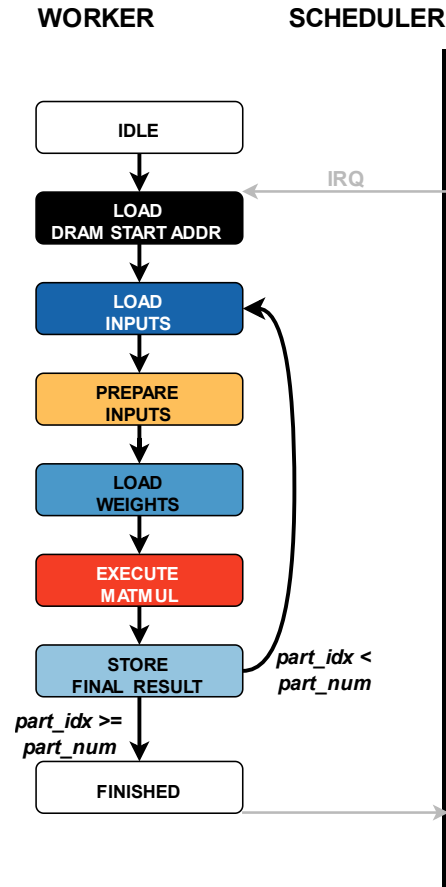
How can we execute DNN layers on SpiNNaker2?

FSM Scheduling on SpiNNaker2

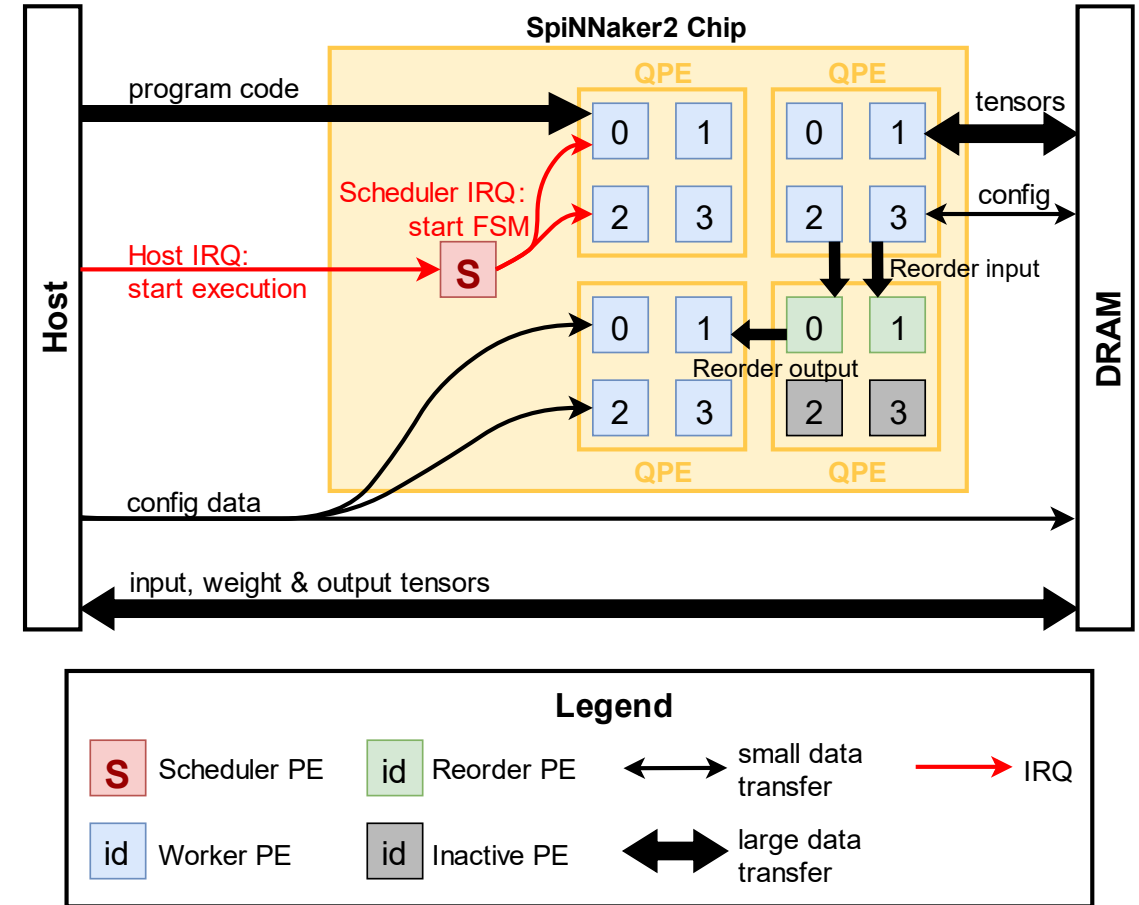
Tiling



FSM



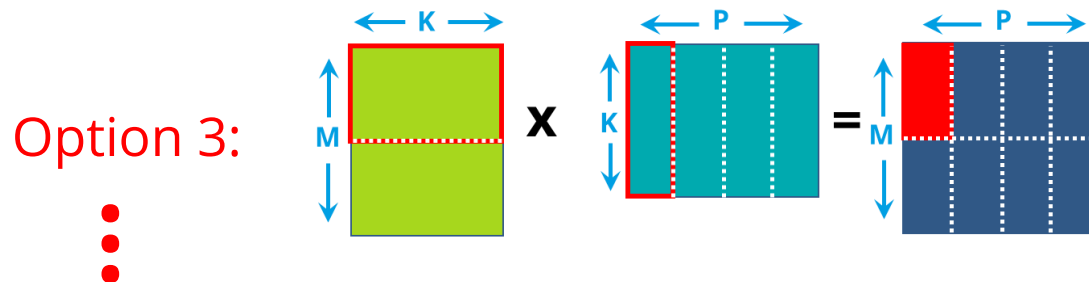
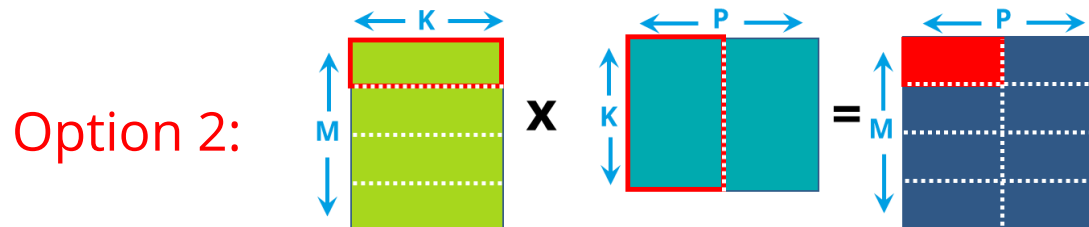
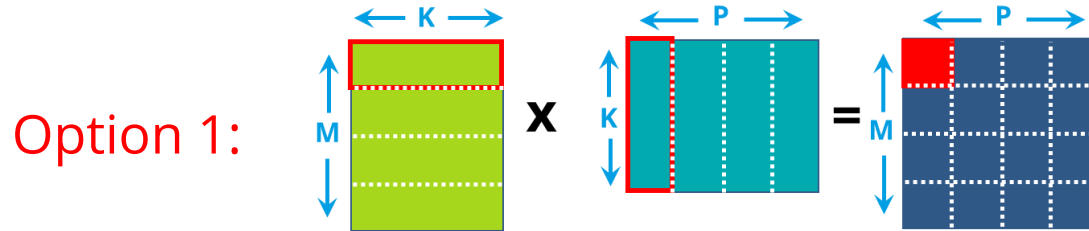
Control & Data Flow



How do we distribute the computation?

Automated Tiling

How to tile a matrix multiplication?



Dimensional Dividers:

$$\mathbb{T}(M) := \{m : (m \mid M) \wedge (m \bmod 4 = 0)\} \subset \mathbb{N}$$

$$\mathbb{T}(P) := \{p : (p \mid P) \wedge (p \bmod 16 = 0)\} \subset \mathbb{N}$$

$$\mathbb{W} := \{w : (w \mid (M \cdot P)) \wedge (w \leq \#(\text{PE})_{\text{avail}})\} \subset \mathbb{N}$$

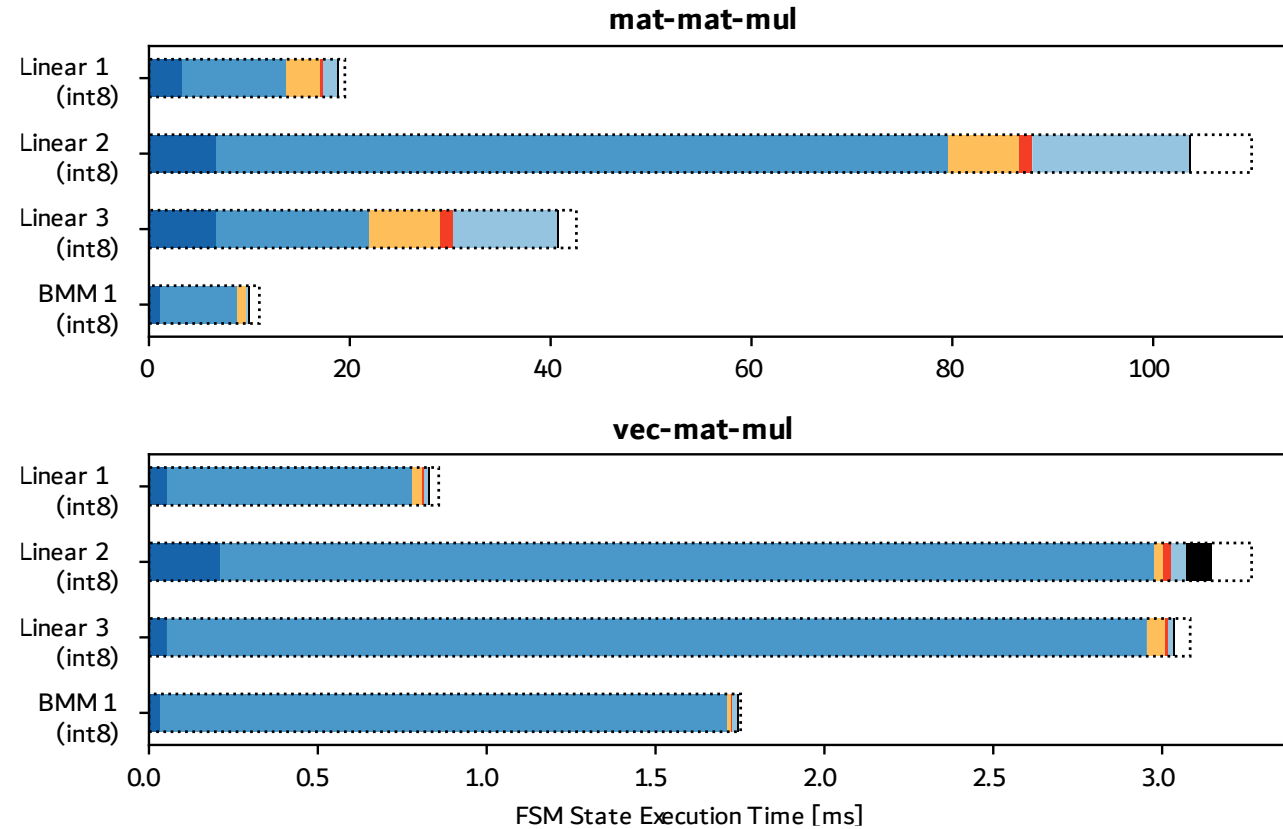
Tiling Space:

$$\mathbb{S}_{\text{MatMul}} = \mathbb{T}(M) \times \mathbb{T}(P) \times \mathbb{W}$$

Priority	Type	Parameter	Condition / Sorting
0	hard	$\#(\text{PE})_{\text{WORKER}}^{(\text{part})}$	$\in \mathbb{N}$
0	hard	$\#(\text{PE})_{\text{WORKER}}^{(\text{part})}$	$\leq \#(\text{PE})_{\text{avail}}$
0	hard	$\text{size}(\text{SRAM})_{\text{max}}^{(\text{PE})}$	$\leq \text{size}(\text{SRAM})_{\text{avail}}^{(\text{PE})}$
0	hard	$\#(\text{parts})$	$\in \mathbb{N}$
1	soft	$\#(\text{parts})$	min
2	soft	$\#(\text{PE})_{\text{WORKER}}^{(\text{part})}$	max

Hard and soft constraints for finding an optimal hardware-aware partitioning for matrix multiplications.

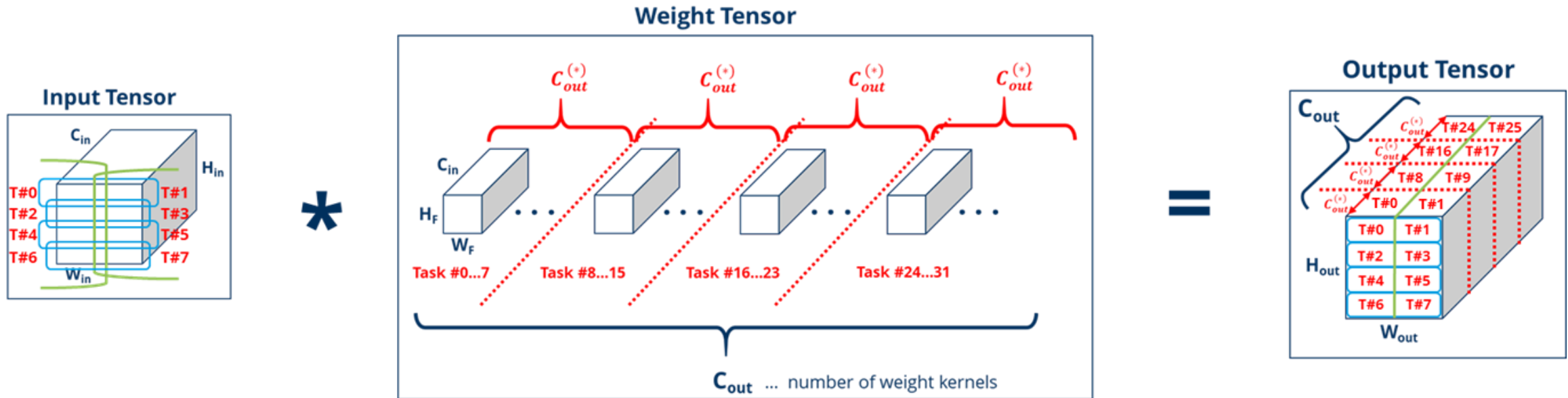
Matrix Multiplication: Large Impact of Data Transfers!



Could we avoid DRAM transfers by handling intermediate activations on-chip?

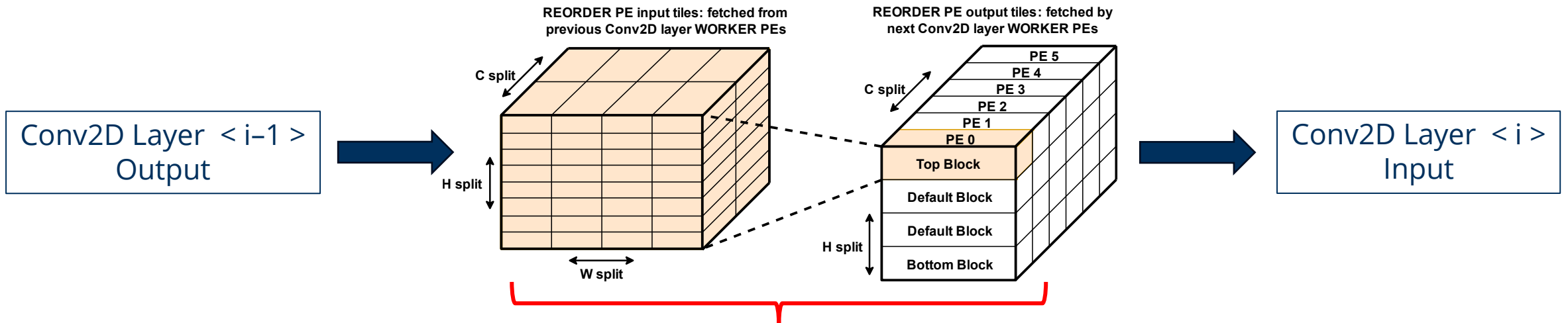
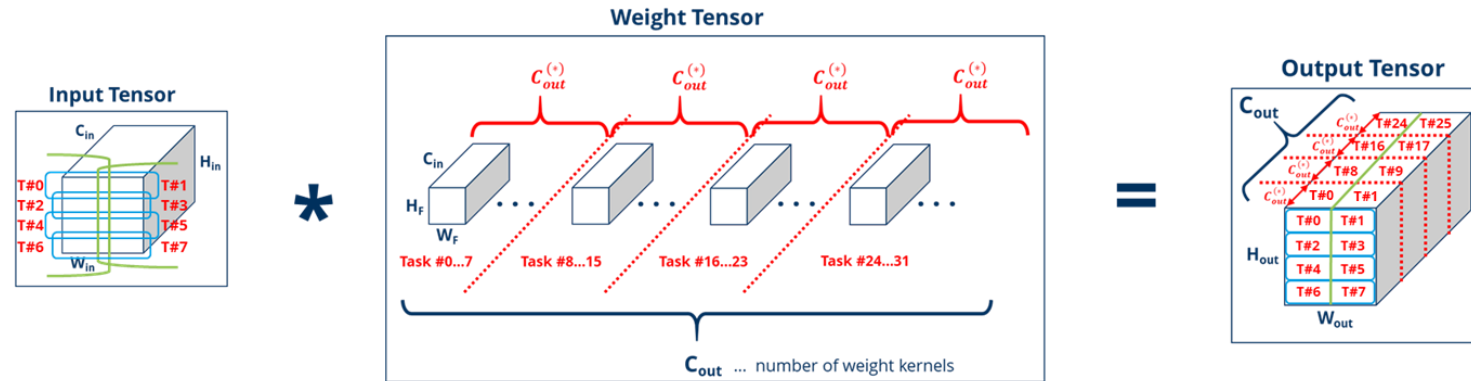
Convolutional Layers

Tiling



Convolutional Layers

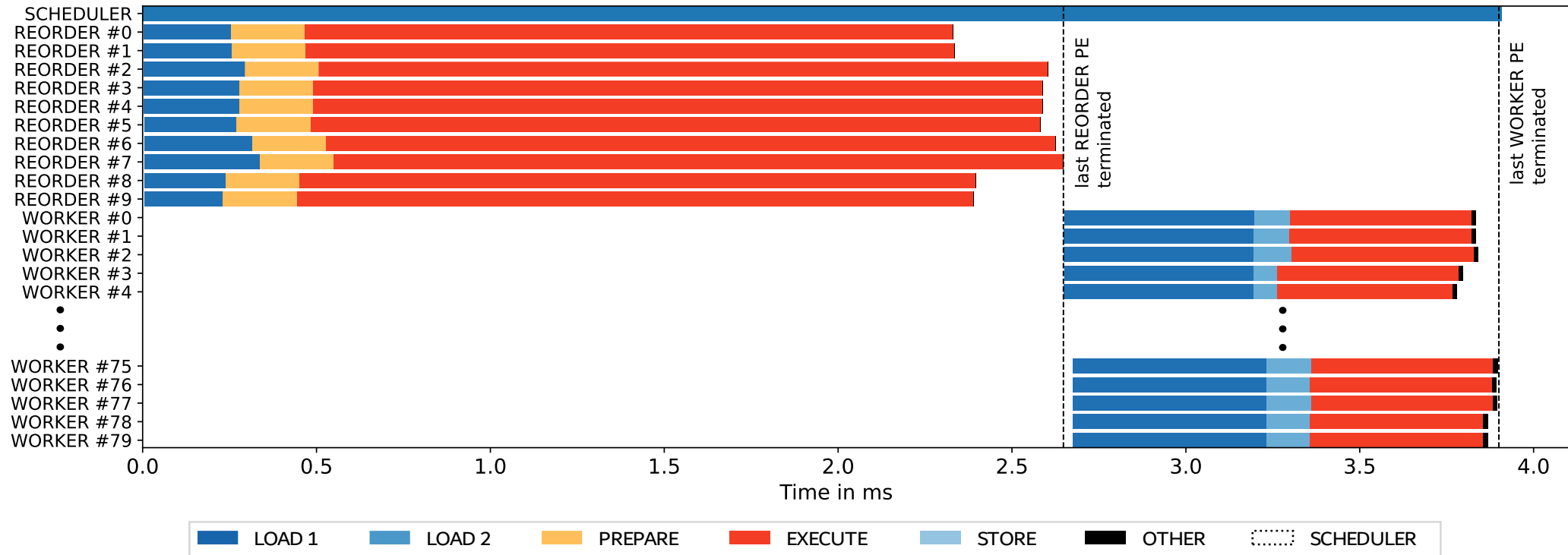
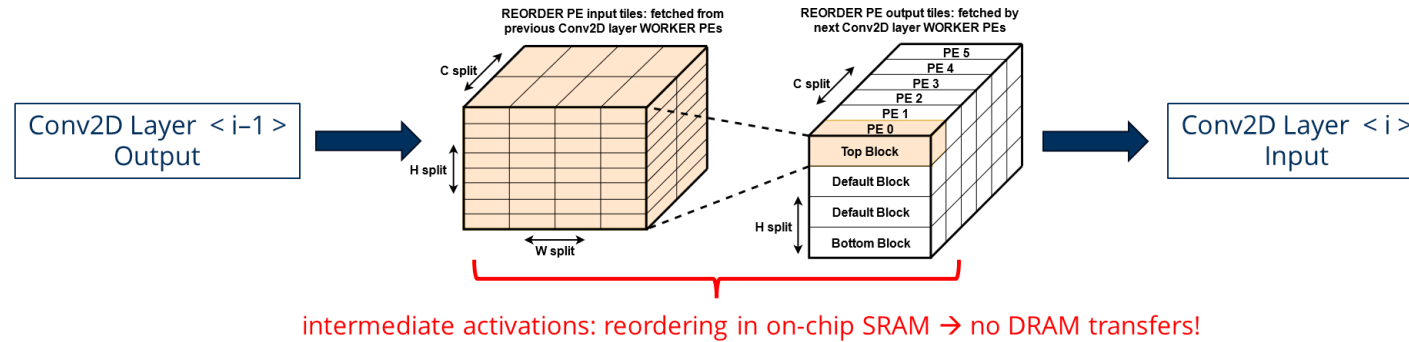
Tiling & Transfer of Intermediate Activations



intermediate activations: reordering in on-chip SRAM → no DRAM transfers!

Convolutional Layers

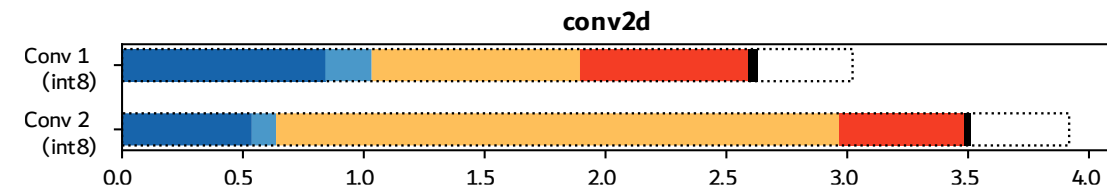
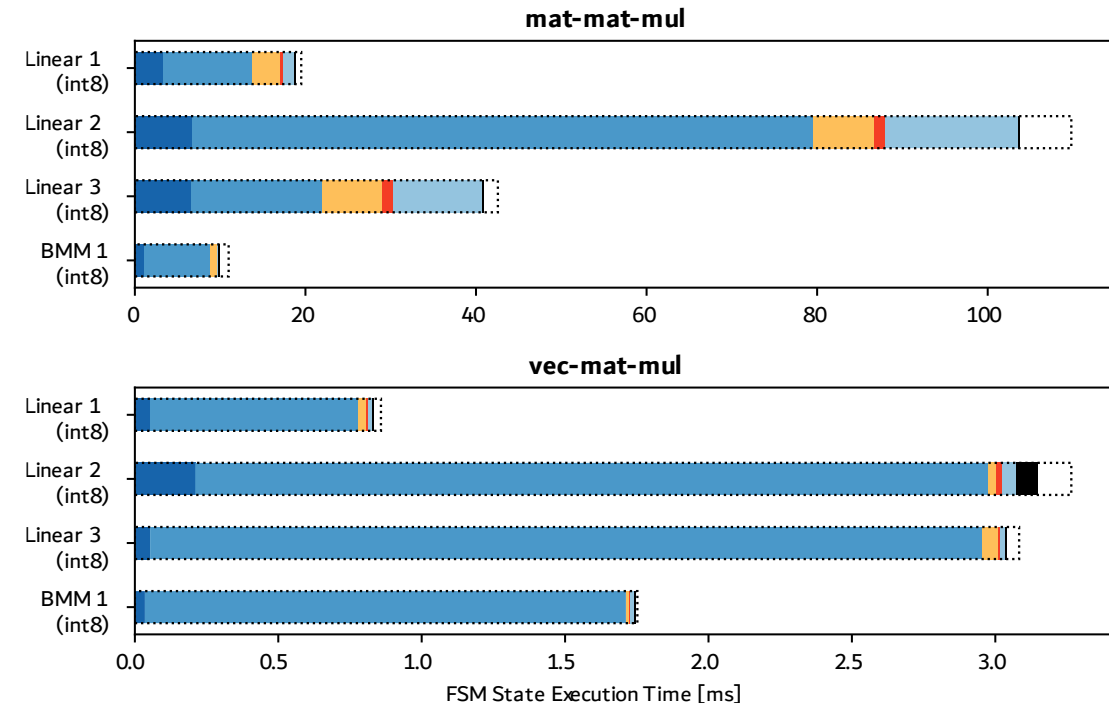
Reordering



Conv2D: Much Less Impact of Data Transfers!

Matrix Multiplication

Conv2D



OctopusScheduler

Summary

What have we developed?

- framework for execution of single DNN layers on SpiNNaker2
- on-chip scheduling for SpiNNaker2
- hardware-aware strategy for automated tiling design space exploration for SpiNNaker2
- on-chip reordering mechanism for Conv2D on SpiNNaker2

Work in progress:

- multi-layer scheduling
- application to CNN & transformer architectures
- optimizations
- pipelining

What has not been covered:

- different other layer types (Softmax, Add, LayerNorm)
- tiling optimization for matrix multiplications
- detailed principle of Conv2D on-chip reordering & padding

→ You can find it in our paper!

OctopusScheduler: On-Chip DNN Scheduling on the SpiNNaker2 Neuromorphic MPSoC

Tim Langer*, Matthias Jobst*[†], Chen Liu*, Florian Kelber*, Bernhard Vogginger*, Christian Mayr*^{†‡}

*Chair of Highly-Parallel VLSI Systems and Neuro-Microelectronics, TU Dresden, Germany

[†]Centre for Tactile Internet with Human-in-the-Loop (CeTI), TU Dresden, Germany [‡]ScaDS.AI Dresden/Leipzig, Germany
{tim.langer, matthias.jobst2, chen.liu, florian.kelber, bernhard.vogginger, christian.mayr}@tu-dresden.de

We would like to thank: