



RISC-V LLM Accelerator Design Platform on C2RTL System Design Verification Framework

Tsuyoshi Isshiki

*Dept. of Information and Communications Engineering
Institute of Science Tokyo*

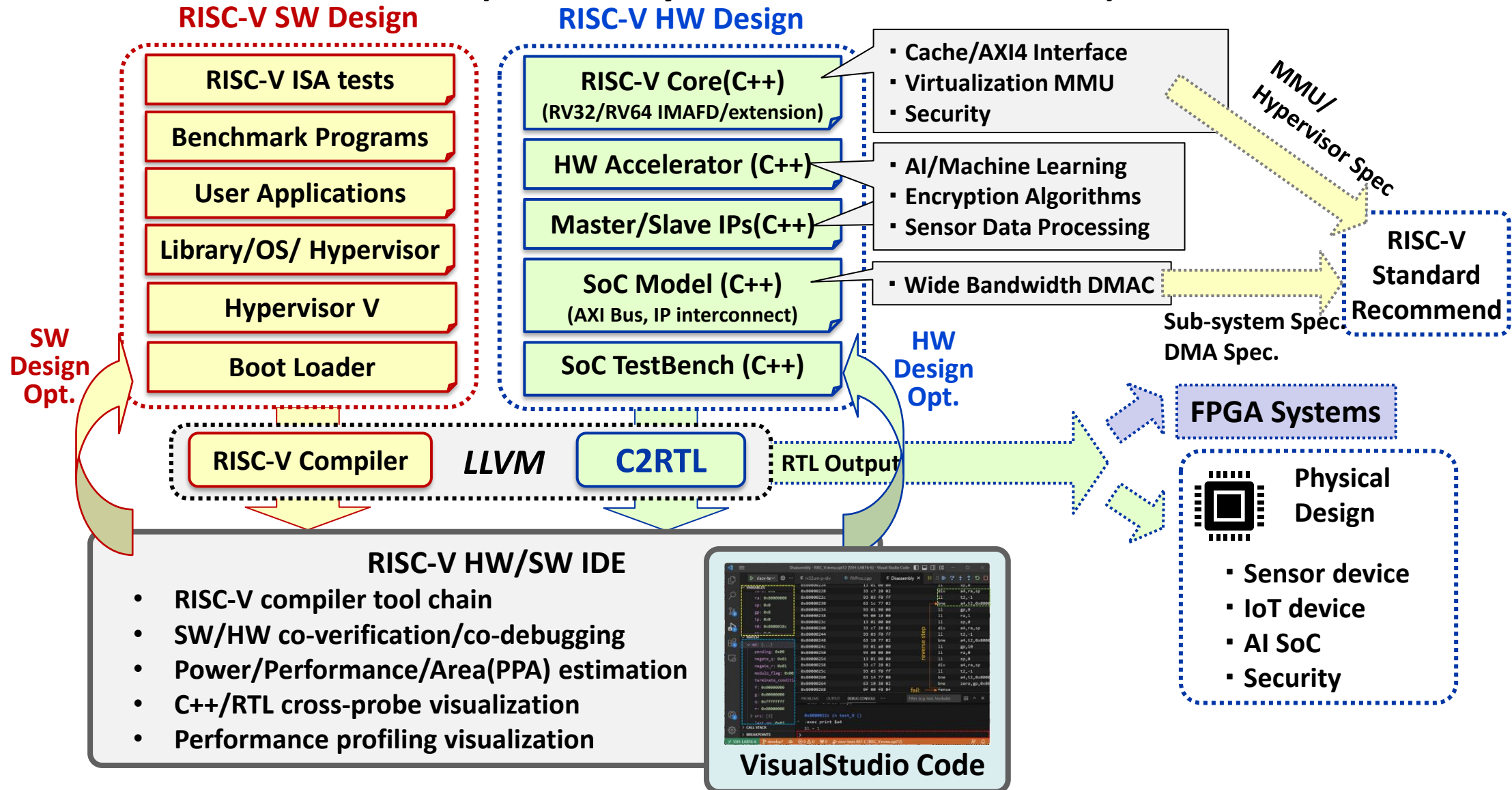
June 23th, 2026

RISC-V Based AI SoCs

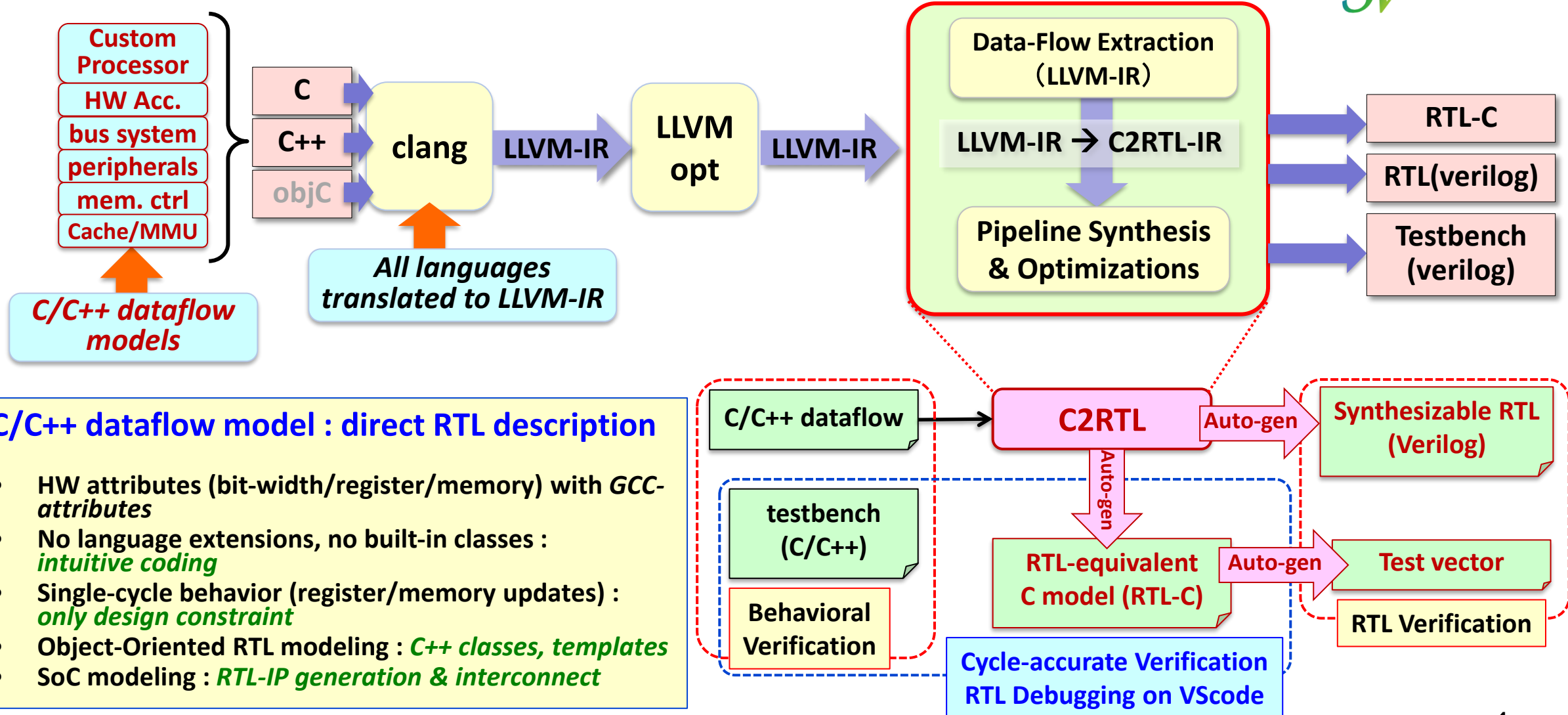
- **25 billion RISC-V Based AI SoCs by 2027 (Semico Research)**
 - RISC-V Based AI SoC revenue : \$291B (2027)
 - RISC-V CPU Semiconductor IP (SIP) 34.9% CAGR through 2027
 - RISC-V SW ecosystem expands adoption in consumer/enterprise/data center/automotive SoC markets
 - RISC-V ratifications : vector, scalar cryptography, hypervisor → RVA23 profile
 - **Design opportunities and challenges of RISC-V architecture**
 - RISC-V ISA specifications help build SW ecosystems for RISC-V based SoCs
 - Large opportunities for drastic improvements in compute/power efficiencies with custom instruction extensions and HW accelerations
- Need System-Level SW/HW Design Environment for efficiently building customized RISC-V cores, HW accelerators and SoC architectures

RISC-V HW/SW Integrated Design Environment (IDE)

(Funded by NEDO: 2022.8 – 2025.3)



LLVM-based C2RTL Framework

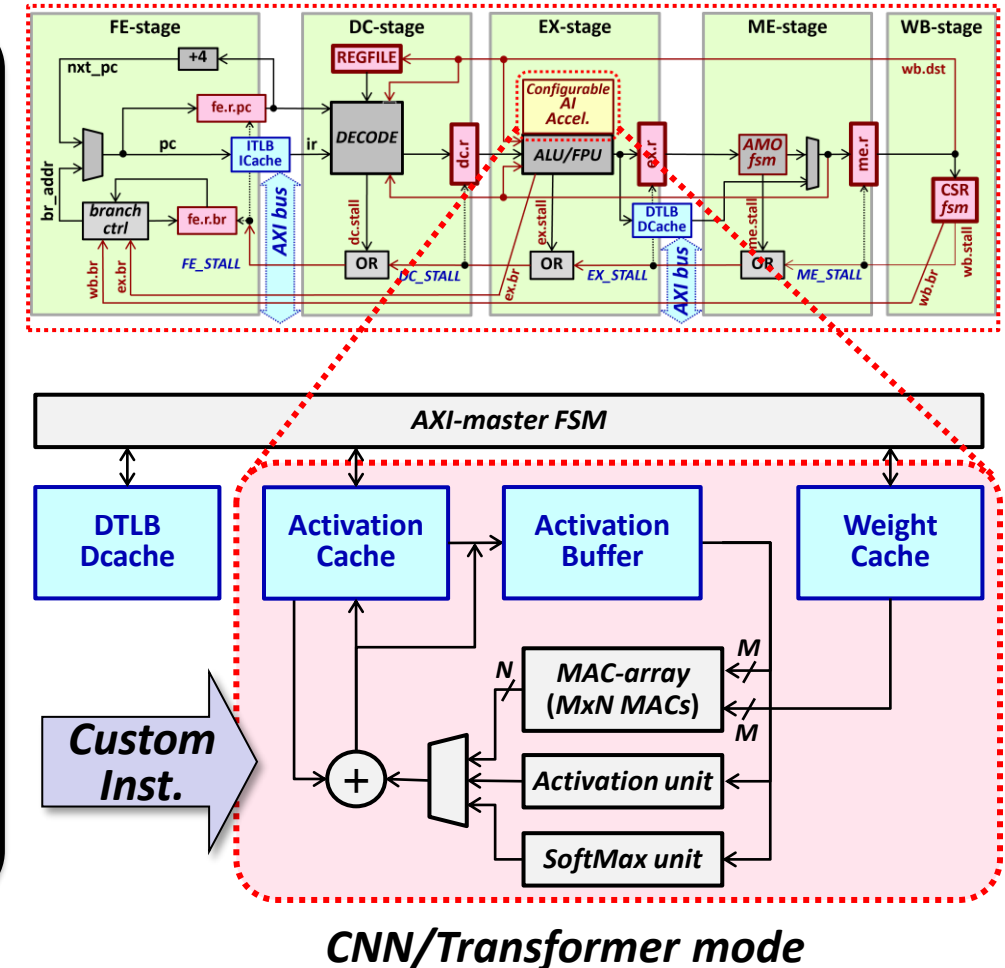


RISC-V AI Accelerator Design Platform

(JST-Research and Development Program for Next-generation Edge AI Semiconductors : 2025 – 2030)

- **Target AI / analytics algorithms**
 - AI : Transformers (LLM, SLM), CNN, etc.
 - Analytics : FFT, wavelet, PCA, etc.
- **Accelerator design on C++/C2RTL → auto-generate RTL model & simulation models**
 - MAC-arrays
 - Non-linear function units (SoftMax, activation)
 - High bandwidth dedicated cache system
 - High bandwidth wide AXI-bus
 - SW-controlled AI accelerator (custom inst.)
- **Enable AI algorithm designers to generate optimized hardware on PPA* design space**

PPA* : Power Performance Area



RISC-V Transformer Accelerator Baseline Architecture

- **Data format : MX-INT8 (32 exponent groups, 8-bit integer quantized data) widely supported in latest GPU/NPUs**
 - Efficient integer-based MAT-MUL compute
 - Output vector renormalization to recover MX-INT8 format
- **SW implementation : tensor-level custom instructions**
 - **LoadVector** : weight matrix, input activation vector
 - **StoreVector** : logits output vectors
 - **Compute** : MAT-MUL, RMSNorm, RoPE, Softmax, SWIGLU
 - **Move** : data transfer between on-chip memories
 - **Config** : datapath setup (src/dst memory pointers, FSM counters)
- **MAC units (currently supports 2 ~ 32 units → final target up to 128 ~ 256 units)**
 - **Weight/KV-cache memory** on each MAC unit
 - **32 x 2 INT8 MACs + normalized accumulator / unit → 2048 MACs/cycle @ 32 units**
 - Element-wise **RMSNorm/RoPE/Softmax/SWIGLU** functions → 1 output / unit
- **High bandwidth on-chip/off-chip memory transfers**
 - 64-byte (512-bits) AXI-bus → **64GB/s @ 1GHz clock**
 - 64-byte wide on-chip memory / MAC unit → **2TB/s, 4TOP/s @ 32 units**

MX-INT8 for Efficient Integer MAT-MUL Compute

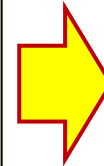
➤ MX-INT8 : 32 exponent groups, 8-bit integer quantized data

$$\begin{aligned} \text{FP vector : } x_{[N]} &= [x_0 \quad x_1 \quad \cdots \quad x_{N-1}] \\ &= \left[2^{x_0^{(E)}} \cdot x_0^{(F)} \quad 2^{x_1^{(E)}} \cdot x_1^{(F)} \quad \cdots \quad 2^{x_{N-1}^{(E)}} \cdot x_{N-1}^{(F)} \right] \end{aligned}$$

Individual exponents on all vector elements

$$\begin{aligned} \text{FP inner-product : } p &= x_{[N]} \cdot y_{[N]} \\ &= \sum_{i=0}^{N-1} 2^{x_i^{(E)} + y_i^{(E)}} \cdot x_i^{(F)} \cdot y_i^{(F)} \end{aligned}$$

Need normalization during accumulation



$$\begin{aligned} \text{MX-INT8 vector : } x_{\langle N \rangle} &= [x_0 \quad x_1 \quad \cdots \quad x_{N-1}] \\ &= 2^{x^{(E)}} [x_0^{(F)} \quad x_1^{(F)} \quad \cdots \quad x_{N-1}^{(F)}] \quad (N = 32) \end{aligned}$$

Single exponent shared by 32 vector elements

$$\begin{aligned} \text{MX-INT8 inner-product : } p &= x_{\langle N \rangle} \cdot y_{\langle N \rangle} \\ &= 2^{x^{(E)} + y^{(E)}} \cdot \sum_{i=0}^{N-1} x_i^{(F)} \cdot y_i^{(F)} \end{aligned}$$

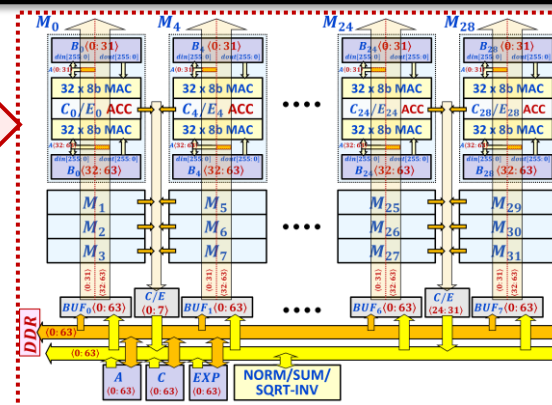
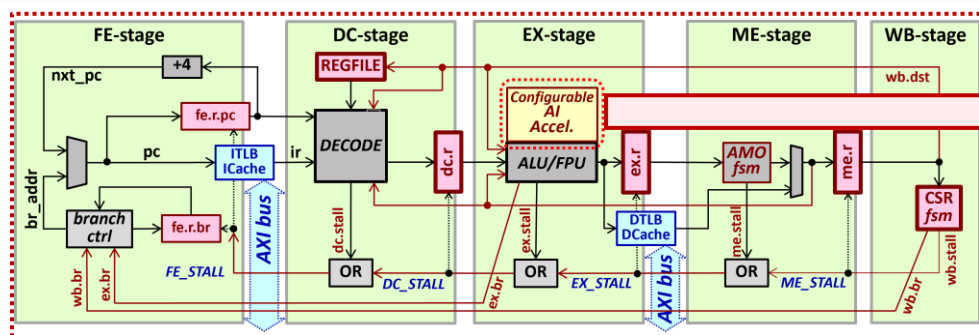
Integer inner-product



Matrix-vector product requires renormalization at each 32 element-group to recover MX-INT8 format

RISC-V Transformer Accelerator Design Flow

- **SW modeling / algorithm optimization** : started from “llama2.c” open-source
 - **MX-INT8 implementation** : exp-grouping strategies, renormalization scheme, special function implementation (Softmax, SWIGLU)
 - Build golden SW model for HW design verification
- **RISC-V Transformer Accelerator HW design verification on C++/C2RTL** :
 - **Tightly-coupled accelerator interface** between RISC-V pipeline micro-architecture
 - **Custom instruction design iteration** : (HW design → SW test) **fully done on C++ !!!**
 - **Gradual architecture refinement** : high-bandwidth off-chip memory transfer (512-bit wide AXI, burst DMA support with cache bypass), on-chip distributed memory, MAC array, data distribution scheme to MAC array, MX-INT8 normalization architecture, etc.



**Transformer HW
Accelerator embedded in
RISC-V pipeline
(in 5000 lines C++ code!!!)**

MX-INT8 Exponent Grouping Scheme on Llama2 Inference

T : Prefill token length, P : token position

$X_{[T],\langle D \rangle}$: input token activation (length : T)

Attention

```

 $X'_{[T],\langle D \rangle} = \text{RMSNorm}(X_{[T],\langle D \rangle}, [W_{\text{RMS\_A}}]_{\langle D \rangle})$ 
for ( $h = 0; h < H_{KV}; h++$ ) { // KV-heads
     $K_{h,[T],\langle D_H \rangle} = \text{RoPE}(X'_{[T],\langle D \rangle} \times [W_K]_{h,[D_H],\langle D \rangle})$ ;
     $KC_{h,[P+T],\langle D_H \rangle} = \text{CONCAT}(KC_{h,[P],\langle D_H \rangle}, K_{h,[T],\langle D_H \rangle})$ ;
     $V_{h,\langle T \rangle,[D_H]} = X'_{[T],\langle D \rangle} \times [W_V]_{h,[D_H],\langle D \rangle}$ ;
     $VC_{h,\langle P+T \rangle,[D_H]} = \text{CONCAT}(VC_{h,\langle P \rangle,[D_H]}, V_{h,\langle T \rangle,[D_H]})$ ;
    for ( $h' = 0; h' < H_Q/H_{KV}; h'++$ ) { // Q-heads
         $\hat{h} = h \cdot H_Q/H_{KV} + h'$ ;
         $Q_{[T],\langle D_H \rangle} = \text{RoPE}(X'_{[T],\langle D \rangle} \times [W_Q]_{\hat{h},[D_H],\langle D \rangle})$ ;
         $S_{[T],\langle P:T \rangle} = \frac{1}{\sqrt{D_H}} Q_{[T],\langle D_H \rangle} \times KC_{h,[P+T],\langle D_H \rangle}$ ;
         $A_{[T],\langle P:T \rangle} = \text{SoftMax}(S_{[T],\langle P:T \rangle})$ ;
         $O_{[T],\langle D_H \rangle} = A_{[T],\langle P:T \rangle} \times VC_{h,\langle P+T \rangle,[D_H]}$ ;
         $X_{[T],\langle D \rangle} += O_{[T],\langle D_H \rangle} \times [W_O]_{\hat{h},\langle D_H \rangle,[D]}$ ;
    }
}
    
```

Weight partitioning scheme for small on-chip memory footprint

$$V_{h,[T],\langle D_H \rangle} = X'_{[T],\langle D \rangle} \times [W_V]_{h,[D_H],\langle D \rangle};$$

→ translates to :

for ($t = 0; t < T; t++$)

for ($k = 0; k < D_H; k++$)

$$V_{h,t,k} = \sum_{d=0}^{D-1} X'_{t,d} \cdot [W_V]_{h,k,d}$$

$K_{h,[T],\langle D_H \rangle}$: exp-group on D_H dim

$V_{h,\langle T \rangle,[D_H]}$: exp-group on T dim

Feed-Forward

```

 $X'_{[T],\langle D \rangle} = \text{RMSNorm}(X_{[T],\langle D \rangle}, [W_{\text{RMS\_F}}]_{\langle D \rangle})$ 
for ( $h = 0; h < H_M; h++$ ) { // M-blocks
     $H1_{[T],\langle M_H \rangle} = X'_{[T],\langle D \rangle} \times [W_1]_{h,[M_H],\langle D \rangle}$ ;
     $H2_{[T],\langle M_H \rangle} = X'_{[T],\langle D \rangle} \times [W_3]_{h,[M_H],\langle D \rangle}$ ;
     $H3_{[T],\langle M_H \rangle} = \text{Swish}(H1_{[T],\langle M_H \rangle}) \odot H2_{[T],\langle M_H \rangle}$ ;
     $X_{[T],\langle D \rangle} += H3_{[T],\langle M_H \rangle} \times [W_2]_{h,[D],\langle M_H \rangle}$ ;
}
    
```

$$S_{[T],\langle P:T \rangle} = Q_{[T],\langle D_H \rangle} \times KC_{h,[P+T],\langle D_H \rangle};$$

→ translates to :

for ($t = 0; t < T; t++$)

$$S_{t,P+t} = \sum_{k=0}^{D_H-1} Q_{t,k} \cdot KC_{h,P+t,k}$$

V-cache exp-group scheme has large impact on design complexity, speed and model accuracy

1% PPL accuracy drop vs FP32 on Llama2-7B

On-chip memory footprint :

Llama2-1.1B : 0.946 MB

Llama2-7B : 2.70 MB

Llama3-8B : 3.78 MB

Tensor-Level Custom Instructions

C Macro call (inline assembler)	ID names : dstID, srcID, compID, moveID, cfgID
<i>Each custom instruction triggers hardwired pipelined FSMs</i> TFM_LoadVector (addr , dstID)	TFM_Weight_Exp, TFM_Logits_Weight_Exp, TFM_VecA_Exp, TFM_VecC_Exp, TFM_LUT_Exp, TFM_Weight_Data, TFM_Logits_Weight_Data, TFM_VecA_Data, TFM_VecC_Data, TFM_LUT_Data, TFM_RMSWeight_Data, TFM_Rope_Data, TFM_Model_Param, TFM_ExtSizeInfo, TFM_MemOffsetInfo, TFM_RMS_Param, TFM_Counter, TFM_ArchInfo
TFM_StoreVector (addr , srcID)	TFM_VecA_Exp, TFM_VecC_Exp, TFM_VecA_Data, TFM_VecC_Data
TFM_Compute (compID)	TFM_COMP_MatMulColNorm, TFM_COMP_MatMulRowNorm, ← <i>For V-cache</i> TFM_COMP_MatMulAcc, TFM_COMP_Rope, TFM_COMP_RMS, TFM_COMP_SoftMax, TFM_COMP_Swig
TFM_Move (moveID)	TFM_MV_C2A, TFM_MV_A2B, TFM_MV_VCache, TFM_MV_VCache_RowExp
TFM_Config (cfgID)	TFM_CFG_RMS, TFM_CFG_KCache, TFM_CFG_KQ, TFM_CFG_VCache, TFM_CFG_VCacheTail, TFM_CFG_V, TFM_CFG_Score, TFM_CFG_AttScore, TFM_CFG_WoAcc, TFM_CFG_H1, TFM_CFG_H3, TFM_CFG_Swig, TFM_CFG_W2Acc, TFM_CFG_FinalRMS, TFM_CFG_Logits
TFM_PosTokenLen (pos , tk_len)	pos : token position, tk_len : input token length

Tensor-Level Custom Instructions

Attention Score Calculation

$$Q_{[T],\langle D_H \rangle} = \text{RoPE} \left(X'_{[T],\langle D \rangle} \times [W_Q]_{h',[D_H],\langle D \rangle} \right);$$

$$S_{[T],\langle P,T \rangle} = \frac{1}{\sqrt{D_H}} Q_{[T],\langle D_H \rangle} \times K_{h,[P+T],\langle D_H \rangle};$$

$$A_{[T],\langle P,T \rangle} = \text{SoftMax}(S_{[T],\langle P,T \rangle});$$

$$O_{[T],\langle D_H \rangle} = A_{[T],\langle P,T \rangle} \times V_{h,\langle P+T \rangle,[D_H]};$$

$$X_{[T],\langle D \rangle} += O_{[T],\langle D_H \rangle} \times [W_O]_{h',\langle D_H \rangle,[D]};$$

```
TFM_Config (TFM_CFG_KQ); // Q projection
TFM_Load (w0[3], TFM_Weight); // LOAD (Wq);
TFM_Compute (TFM_COMP_MatMulColNorm); // Q ← X' * Wq;
TFM_Compute (TFM_COMP_Rope); // Q ← Rope(Q);
TFM_Config (TFM_CFG_Score); // SoftMax score
TFM_Compute (TFM_COMP_MatMulColNorm); // S ← Q * K;
TFM_Compute (TFM_COMP_SoftMax); // A ← SoftMax(S);
TFM_Config (TFM_CFG_AttScore); // att_score
TFM_Compute (TFM_COMP_MatMulColNorm); // O ← A * V;
TFM_Move (TFM_MV_C2A); // relocate O : Cmem → Amem
TFM_Config (TFM_CFG_AttAcc); // att_output
TFM_Load (w0[4], TFM_Weight); // LOAD (Wo);
TFM_Compute (TFM_COMP_MatMulAcc); // X ← X + O * Wo;
```

```
#define TFM_Load(qv, dtype) TFM_LoadVector(qv.e, dtype##_Exp); TFM_LoadVector(qv.q, dtype##_Data)
```

```
TFM_Load (w0[3], TFM_Weight);
```

```
TFM_LoadVector (w0[3].e, TFM_Weight_Exp);
TFM_LoadVector (w0[3].q, TFM_Weight_Data);
```

Tensor-Level Custom Instructions

Feed-Forward

```
for (h = h0; h < HM; h++) {
  H1[T],⟨MH⟩ = X'[T],⟨D⟩ × [W1]h,⟨D⟩,[MH] ;
  H2[T],⟨MH⟩ = X'[T],⟨D⟩ × [W3]h,⟨D⟩,[MH] ;
  H3[T],⟨MH⟩ = Swish(H1[T],⟨MH⟩) ⊙ H2[T],⟨MH⟩ ;
  X[T],⟨D⟩ += H3[T],⟨MH⟩ × [W2]h,[D]⟨MH⟩ ;
}
```

```
#define TFM_Load (qv, dtype) ¥
TFM_LoadVector (qv.e, dtype##_Exp); ¥
TFM_LoadVector (qv.q, dtype##_Data)
```

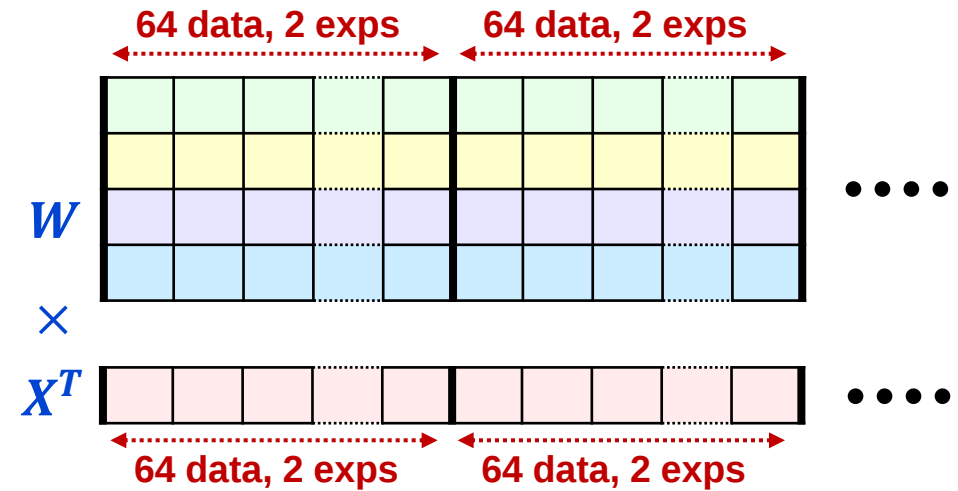
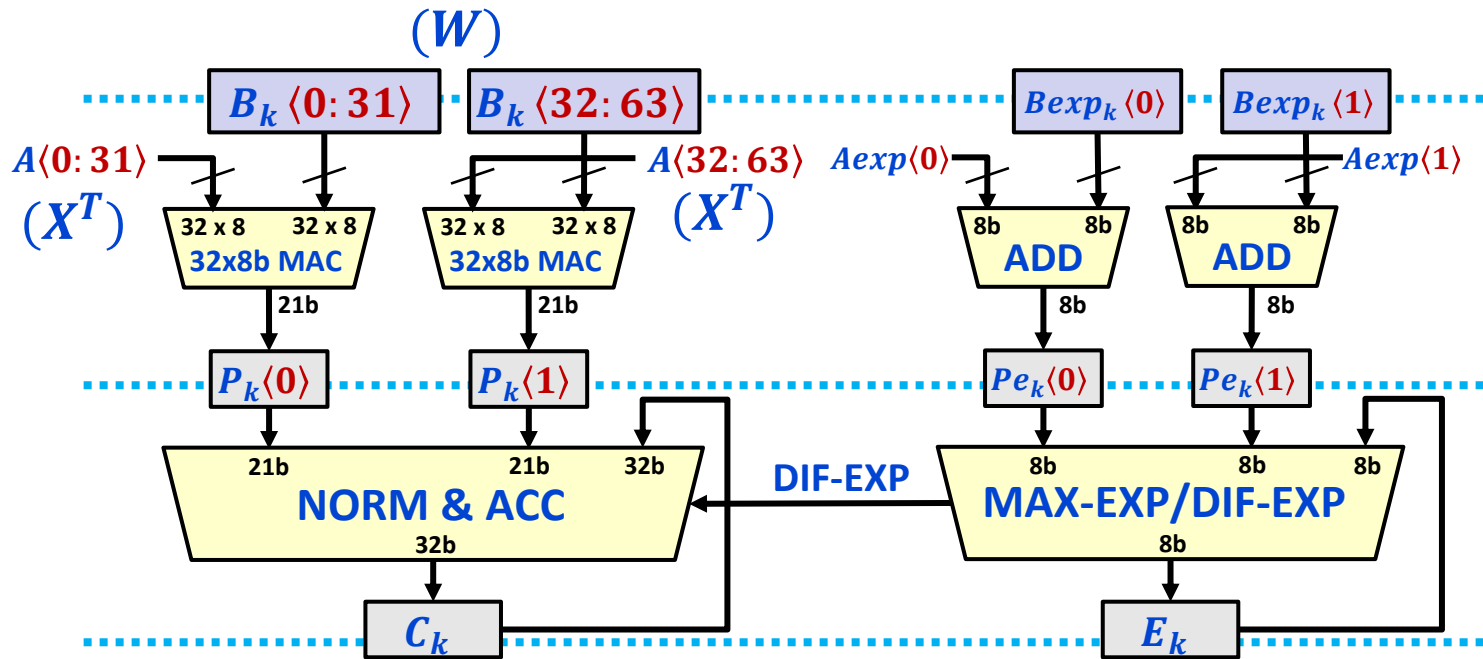
```
TFM_Load (w0[6], TFM_Weight);
```

```
for (int h = 0; h < model_param.num_hblks; h++) {
  TFM_Config (TFM_CFG_H1); // H1 projection
  TFM_Load (w0[6], TFM_Weight); // LOAD (W1);
  TFM_Compute (TFM_COMP_MatMulColNorm); // H1 ← X' x W1;
  TFM_Config (TFM_CFG_H3); // H2 projection
  TFM_Load (w0[8], TFM_Weight); // LOAD (W3);
  TFM_Compute (TFM_COMP_MatMulColNorm); // H2 ← X' x W3;
  TFM_Config (TFM_CFG_Swig); // SWIGLU
  TFM_Compute (TFM_COMP_Swig); // H3 ← Swish(H1) * H2;
  int last_hblk = (h == model_param.num_hblks - 1);
  TFM_Config2 (TFM_CFG_SwigAcc, (last_hblk << 1)); // FFN output
  TFM_Load (w0[7], TFM_Weight); // LOAD (W2);
  TFM_Compute (TFM_COMP_MatMulAcc); // X ← X + H3 * W2 ;
  for (int j = 6; j < 9; ++j) { // increment weight pointers
    QVec_inc_pos(&w0[j], QVEC_INC_DEPTH(wqvi[j], last_hblk));
  }
}
```

```
TFM_LoadVector (w0[6].e, TFM_Weight_Exp);
TFM_LoadVector (w0[6].q, TFM_Weight_Data);
```

MAC Unit : MAC Pipeline

- **MAC units (2 ~ 32 units) → 2048 MACs/cycle**
 - Weight/KV-cache memory
 - 32 x 2 INT8 MACs** + normalized accumulator
 - 2-stage pipeline : MAC-stage → NORM/ACC-stage



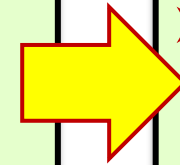
W : stored at MAC unit memory
 X^T : broadcast to all units

MAC pipeline : 64 MACs / cycle at each MAC unit

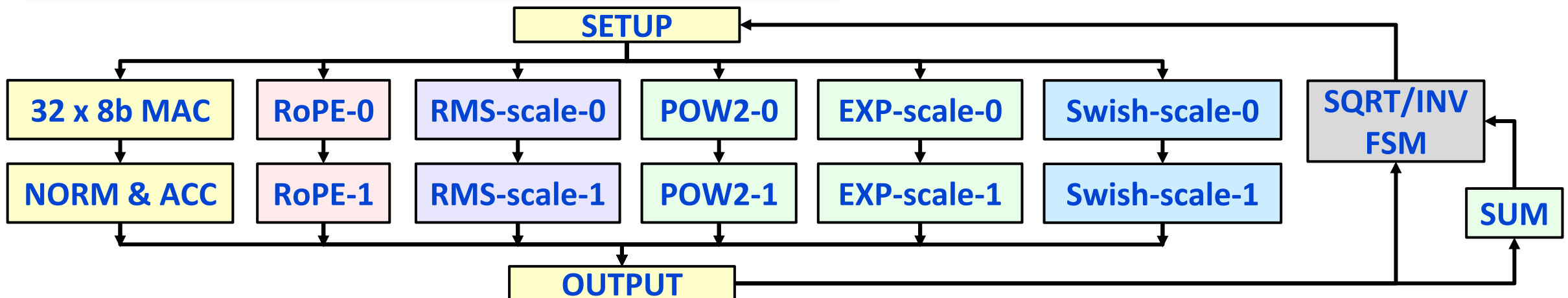
MAC Unit : Element-Wise Special Function Pipelines

- **RoPE** : Pairwise rotation on K/Q activations
- **RMSNorm** :
 - Squared sum (reuse MAC pipeline : $A \cdot A$)
 - Square-root inverse
 - Scaling with weights
- **Softmax** :
 - $\exp(x) = \text{pow}(2, x \cdot \log_2 e)$
 - Exp-sum inverse
 - Exponent sum-inverse scaling
- **SWIGLU** : Element-wise scaling of Swish(H1) and H2

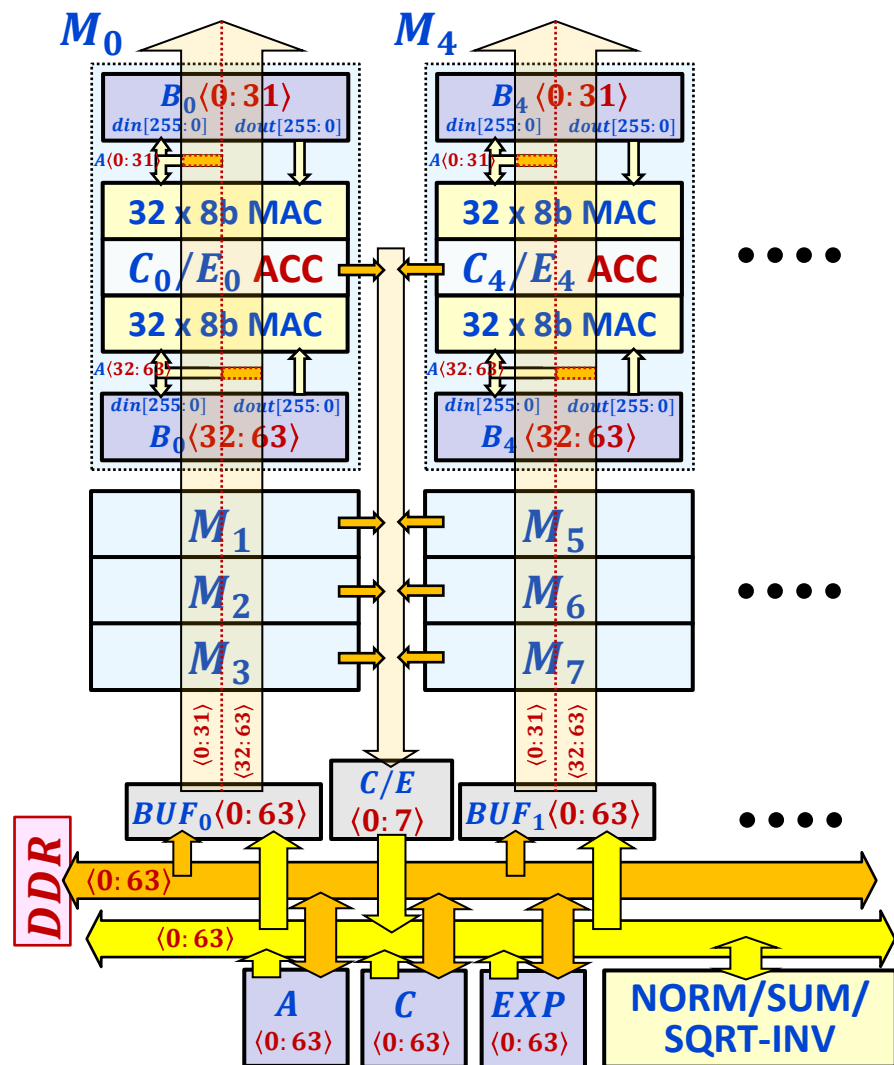
- **PWL LUTs** :
 - $\text{pow}(2, x)$
 - $\text{Swish}(x) = x / (1 + \exp(x))$
- **SQRT/INV FSM (Newton-Raphson)** :
 - Square-root inverse : $x^{-0.5}$
 - Inverse : x^{-1}



All special function pipelines align with the MAC pipeline



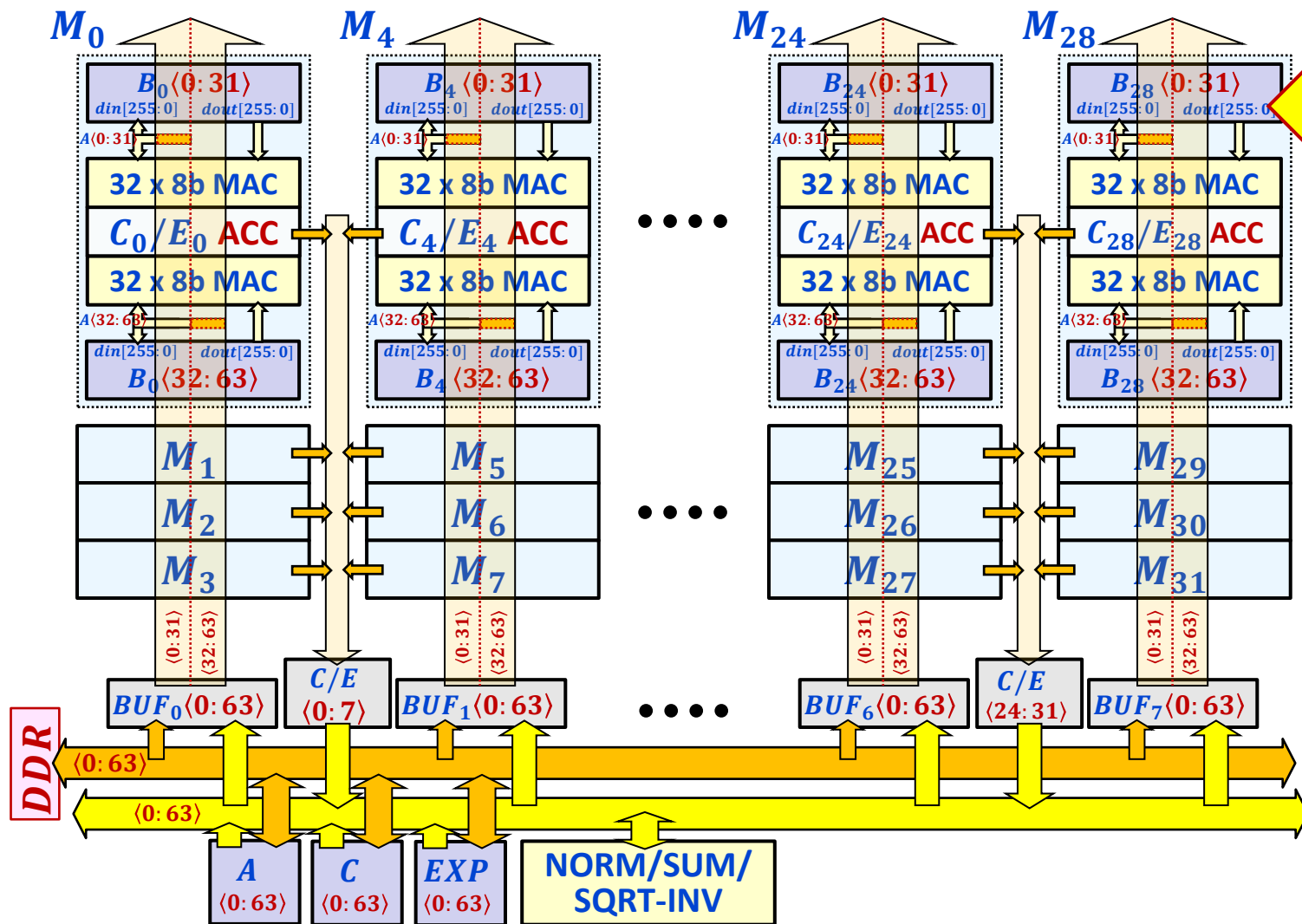
Full MX-INT8 Pipeline with Output Normalization



$M_{R,C} \cdot X_{C,T}$ matrix multiplication : $R \cdot C \cdot T$ MACs
 $\rightarrow (R \cdot C \cdot T)/N + (7 + 32/N)$ cycles

stage	Main pipeline operations (N : # of MAC units)
0	Read Amem (64 activations)
1	Broadcast 64 activations
2	$32 \times 8b \times N$ MACs / element-wise FUNC(0)
3	Normalize & accumulate / element-wise FUNC(1)
4	Latch MAC outputs
5, 6, 7	Column-wise group normalization (If $N < 32$, $32/N - 1$ extra latencies added)
stage	Row-wise group normalization pipeline (VCache)
0, 1	repeats $C \cdot T/N$ times after main pipeline completion

High-Bandwidth On-Chip/Off-Chip Memory Transfer



***Transformer HW Accelerator
embedded in RISC-V pipeline
(in 5000 lines C++ code!!!)***

64-byte wide MAC-unit memory
→ 2TB/s @ 32 units, 1GHz clock

**Decode throughput :
(off-chip bandwidth limited)**

Llama2-1.1B	56.4 tokens/s
Llama2-7B	9.2 tokens/s
Llama3-8B	7.72 tokens/s

64-byte (512-bits) AXI-bus
→ 64GB/s @ 1GHz clock

- DDR4 : 25.6GB/s per channel
- DDR5 : 67.4GB/s per channel
- HBM4 : >2TB/s per stack

Summary

- **RISC-V Transformer Accelerator HW design verification on C++/C2RTL :**

- **Tightly-coupled accelerator interface** between RISC-V pipeline
- **MX-INT8 data format** for efficient integer MAT-MUL compute
- **Tensor-level custom instructions** : LoadVector, StoreVector, Compute, Move, Config
- **2048 MACs/cycle @ 32 MAC units (32 x 2 INT8 MACs /unit)**
- **RMSNorm/RoPE/Softmax/SWIGLU** functions inside MAC unit
- **Off-chip data transfer : 64GB/s @ 1GHz clock**
- **On-chip memory bandwidth : 2TB/s @ 32 MAC units, 4 TOP/s @ 1GHz**
- **Accelerator design in C++ (5000 lines) → automatic RTL generation**

On-chip memory :

Llama2-1.1B : 0.946 MB
Llama2-7B : 2.70 MB
Llama3-8B : 3.78 MB

**Decode throughput :
(off-chip bandwidth limited)**

Llama2-1.1B : **56.4 tokens/s**
Llama2-7B : **9.2 tokens/s**
Llama3-8B : **7.72 tokens/s**

- **On-going works :**

- **FPGA implementation** : most custom instructions tested on small testbench, will soon deploy a full stack LLM demo on Linux
- **Chip implementation** : 32~64 MAC-unit test chip on 16nm (1st tape-out this year)
- **System-level AI accelerator simulator** : PPA analysis, accuracy (perplexity) analysis, architecture exploration on SOTA LLM models for edge AI devices
- **Automatic SW code generation flow from open-source LLM Frameworks (ex. Hugging Face)**

This work was supported by JST-Research and Development Program for Next-generation Edge AI Semiconductors Grant Number JPMJES2515.