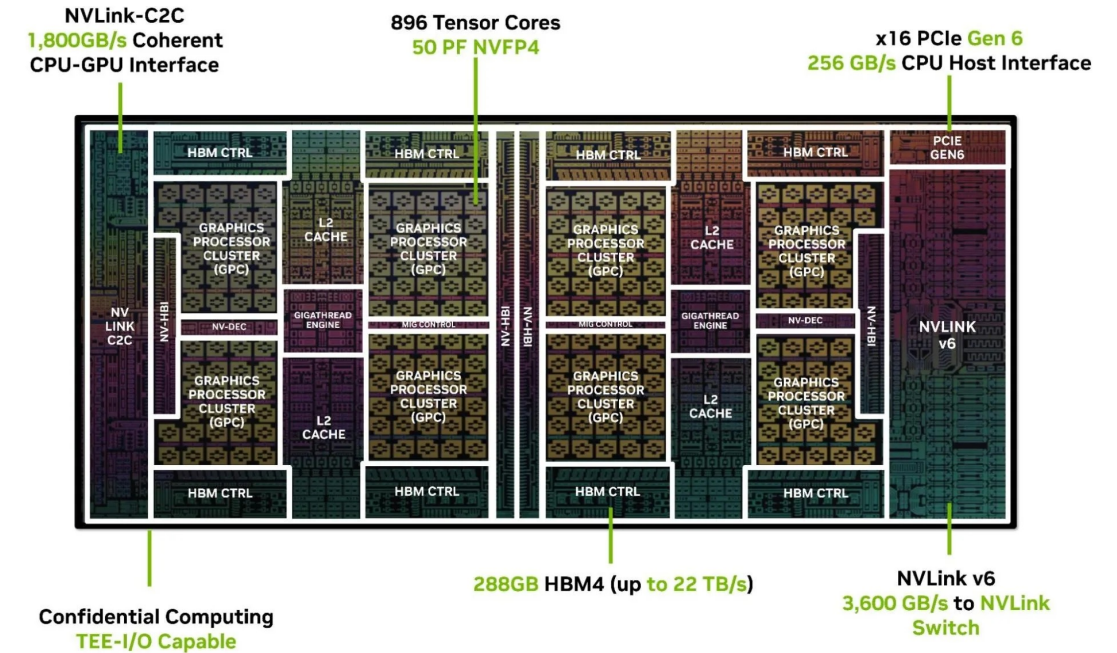


GPU architecture and workload evolution

Workload changes → bottleneck changes → architecture changes

UTILIZING THE ARITHMETIC UNITS

The useful question is whether those units can be fully utilized and which key factors determine their utilization.



Rubin GPUs contain up to 224 SMs and 288GB HBM4 Memory. Available SM count and HBM varies by SKU.

Six design decisions that determine GPU utilization

02

Each decision removes a different reason for execution units to wait.



1 Provide enough parallel work

Many SMs process independent thread groups at the same time.

2 Keep warp state on chip

Keep waiting warps in registers so ready warps can run.

3 Keep reused data nearby

Registers, shared memory, and caches reduce repeated HBM accesses.

4 Match hardware to operations

CUDA cores, Tensor Cores, SFUs, and load/store units use different data paths.

5 Overlap movement and execution

Copy future inputs while current arithmetic runs. Wait before using them.

6 Scale beyond one die

Die links and NVLink carry data needed by work on other processors.

WORKLOAD MILESTONE I

Computer Graphics

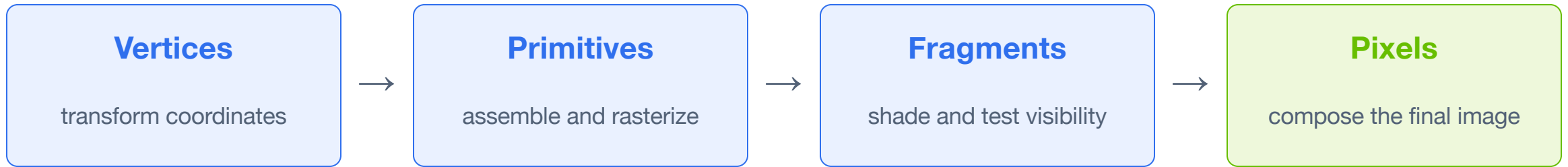
How can we render a complex scene within one frame deadline?

Workload	Render many vertices and fragments with similar operations.
Bottleneck	The machine must finish a large amount of parallel work per second.
Architectural response	Replicate arithmetic units and specialize the rendering stages.

Graphics rewards parallel throughput

04

Many vertices and fragments need similar calculations within one frame deadline.



Parallel work

Many scene elements can be processed at the same time.

Throughput pressure

At 60 frames/s, the entire frame has about 16.7 ms.

Hardware response

Parallel arithmetic and specialized rendering stages increase total work/s.

WORKLOAD MILESTONE II

Scientific and General-Purpose Computing

How can the same parallel hardware run numerical programs directly?

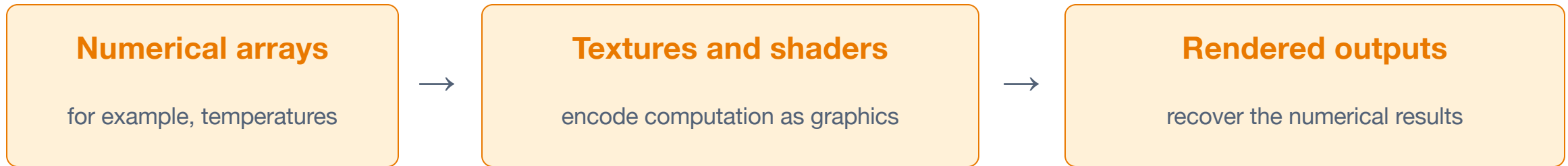
Workload	Simulations and numerical methods update large arrays in parallel.
Bottleneck	Graphics interfaces restrict general data access and cooperation.
Architectural response	Expose programmable threads, on-chip storage, and synchronization.

Scientific computing needed a general programming model

06

The arithmetic fit the GPU, but numerical arrays had to pass through graphics interfaces.

EARLY GPGPU



CUDA



Workload

Simulations and numerical methods expose parallel updates and reusable data.

Programming bottleneck

Graphics APIs make general data access and cooperation awkward.

Hardware and software

CUDA exposes threads, local storage, and synchronization directly.

Why CUDA needed registers and shared memory

07

General programs need persistent values, flexible memory access, and synchronization.

1999

GeForce 256

Specialized graphics stages processed vertices and pixels. General algorithms had to work within graphics memory and execution interfaces.

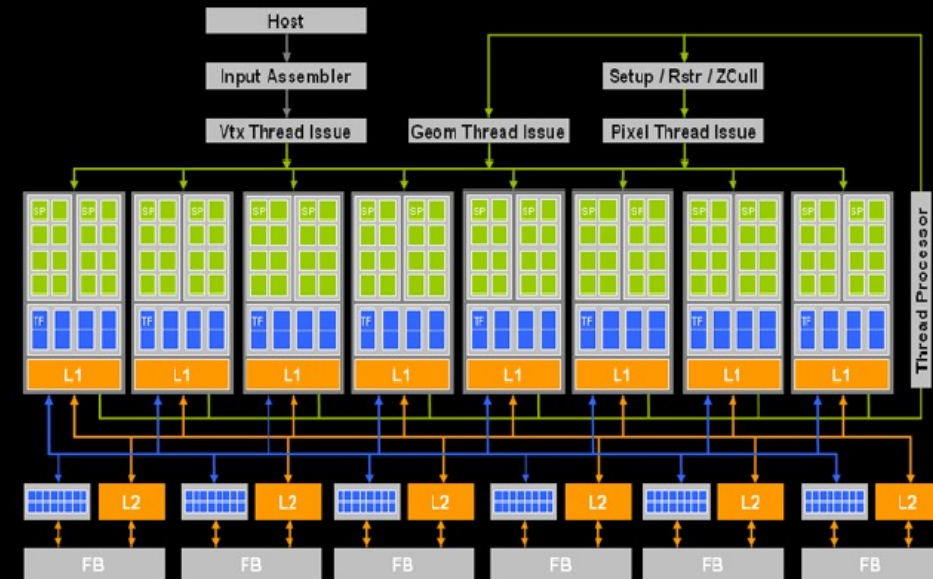
2006

GeForce 8800 / G80

CUDA mapped general programs onto unified processors. G80 combined 128 stream processors with on-chip registers, 16 KB of shared memory per SM, and warp scheduling to reuse data and tolerate off-chip latency.

GeForce 8800 replaces the pipeline model

- The future of GPUs is programmable processing
- So – build the architecture around the processor



GeForce 8800 GTX anchor: 128 stream processors, about 0.52 TFLOPS, 768 MB GDDR3, and 86.4 GB/s memory bandwidth

A CUDA thread: private state, shared execution

A thread runs one instance of a GPU function (kernel). A warp groups 32 threads for instruction issue.

What belongs to one thread?

Identity	threadIdx identifies its position within a block.
Control state	Program counter: next instruction. Call stack: function-call context. Per-thread control state since Volta.
Private values	Registers hold operands and results. Local memory holds other private data and can hold spills.

One instruction, separate register values

Illustrative instruction: $R2 = R0 + R1$

Thread	R0	R1	R2 after
0	2	3	5
1	7	4	11

R0 in thread 0 and R0 in thread 1 name different values. Physical register capacity is shared, but each thread has its own allocation.

How resident warps make progress

Stalled Next instruction must wait.	Eligible Ready, awaiting selection.	Selected Issues an instruction, then continues.
--	--	--

Issue is not completion. Registers stay allocated while waiting. Sources: CUDA Programming Guide and Nsight Compute.

Threads and warps

08

A warp groups 32 threads for instruction scheduling.

PROGRAMMING VIEW

THREAD

One execution of kernel code

A thread has its own ID, registers, and control state.

BLOCK

A cooperating thread group

A block runs on one SM. Its threads can share memory and synchronize.

WARP

The hardware scheduling group

The SM partitions consecutive block threads into groups of 32.

CONCRETE EXAMPLE

One block contains 64 threads

Hardware creates two warps from consecutive thread IDs.



INSTRUCTION ISSUE

The scheduler selects a ready warp. One instruction applies to the warp's active threads.

BRANCH DIVERGENCE

When threads choose different paths, the warp executes each path with a different active mask. Some lanes remain idle.

Following an instruction through the SM

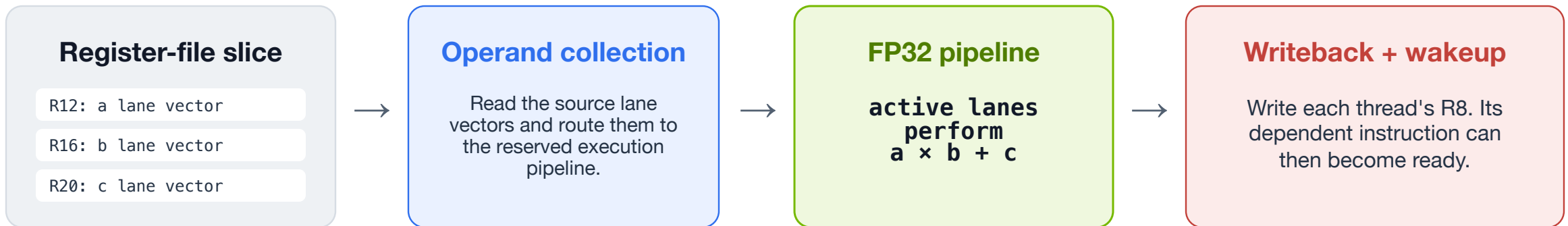
13

Fetch and issue select the operation, then registers supply the data.

CONTROL PATH



OPERAND AND RESULT PATH



LATENCY AND THROUGHPUT

One instruction takes time to complete. Independent instructions can overlap in the pipeline. This is a conceptual data path, not a physical stage or register-port specification.

CUDA cores perform one operation per active thread

12

The warp shares an instruction while each thread supplies its own operands.

ISSUED WARP INSTRUCTION

FFMA R8, R12, R16, R20

$R8 = R12 \times R16 + R20$

REGISTER OPERANDS

lane 0: a_0, b_0, c_0

lane 1: a_1, b_1, c_1

lane 2: a_2, b_2, c_2

...

lane 31: a_{31}, b_{31}, c_{31}

ONE BLACKWELL ULTRA SM PARTITION

32 warp lanes

All active threads receive the operation. A false per-thread condition prevents that thread from writing its result.



Green lanes write results. Gray lanes keep their previous values (lanes 3 and 30 shown).

PIPELINE RESULT

Execute and write back

Each active lane computes its own result and writes it to that thread's destination register.

WHAT LIMITS THROUGHPUT

Dependent instructions wait for results. A saturated arithmetic pipe also delays issue. The scheduler hides these delays by selecting another eligible warp.

Blackwell Ultra: four SM partitions × 32 CUDA cores = 128 CUDA cores per SM. Counts and supported instruction rates vary by architecture.

How a warp becomes ready to issue

10

Resident warps retain their state while dependencies delay their next instruction.

RESIDENT WARP POOL



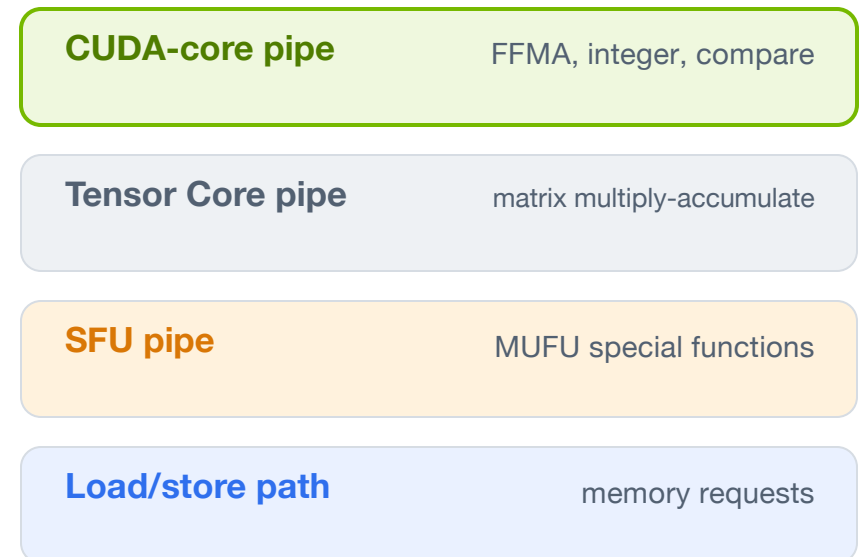
WARP SCHEDULER

A warp can issue when:

- 1 its next instruction is available
- 2 the required operands are ready
- 3 the target pipeline can accept work

Selected this cycle: Warp A

DISPATCH TARGET



State stays on chip

Selecting a resident warp requires no operating-system register save or restore.

Occupancy measures residency

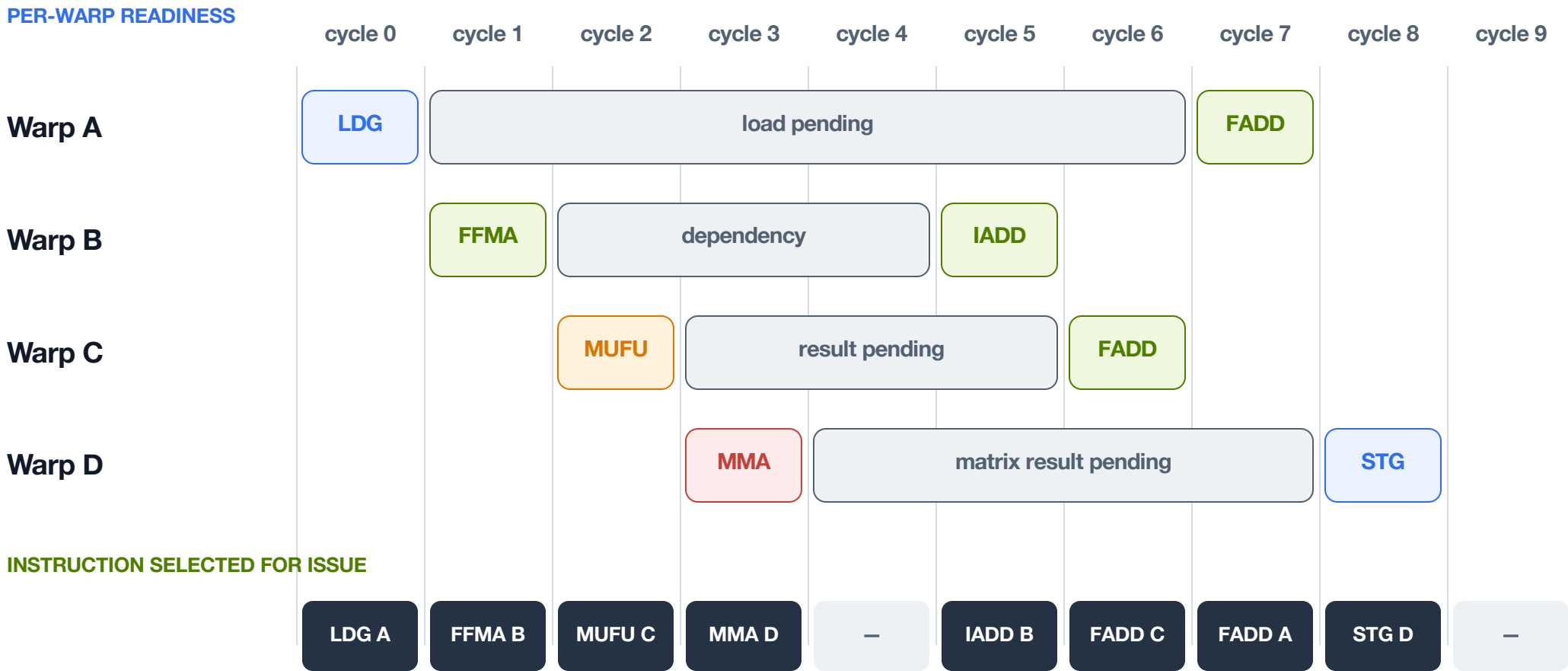
Active resident warps / maximum resident warps. More warps help only if work is ready.

The scoreboard tracks dependencies

Hardware records unfinished results. A dependency, barrier, or unavailable pipe keeps a warp from issuing.

Instruction issue and completion happen at different times

Other warps can issue while an earlier instruction is still in flight.



Latency hiding fills issue slots with independent work while earlier instructions wait.

Illustrative one-issue-per-cycle schedule. Waiting intervals are not measured instruction latencies.

Inside a Blackwell Ultra streaming multiprocessor

09

Modern example: Blackwell Ultra shows where threads and warps execute.



An SM is the complete processing block

A CUDA core is one kind of execution resource inside the SM.

1

Warp schedulers choose ready work

Four schedulers manage assigned warps. Issuing starts work, which can finish later.

2

Registers supply each thread's operands

Threads share physical capacity, but each thread retains its own register values.

3

Specialized units execute different instructions

CUDA cores handle general arithmetic. Tensor Cores handle matrix operations. SFUs handle functions such as exponentials.

4

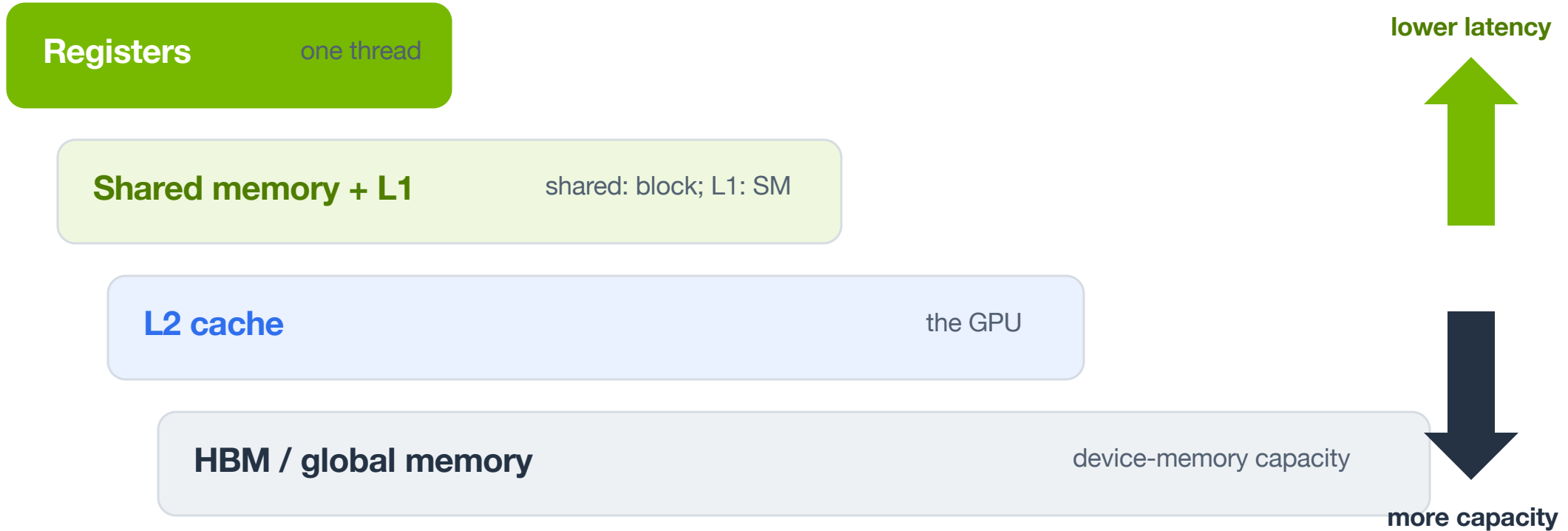
On-chip storage feeds the units

Shared memory, caches, registers, and Tensor Memory reduce repeated traffic to HBM.

Where GPU data live

14

Reusing nearby data reduces traffic to larger, more distant storage.



Reuse: load a tile from HBM, retain it on chip, and use it many times.

Fermi added a general cache hierarchy

15

Historical response: Fermi added general caching in 2010.

2010 • FERMI

General-purpose caching

Each SM gained an L1 cache. A unified L2 cache served load, store, and texture requests across the GPU.

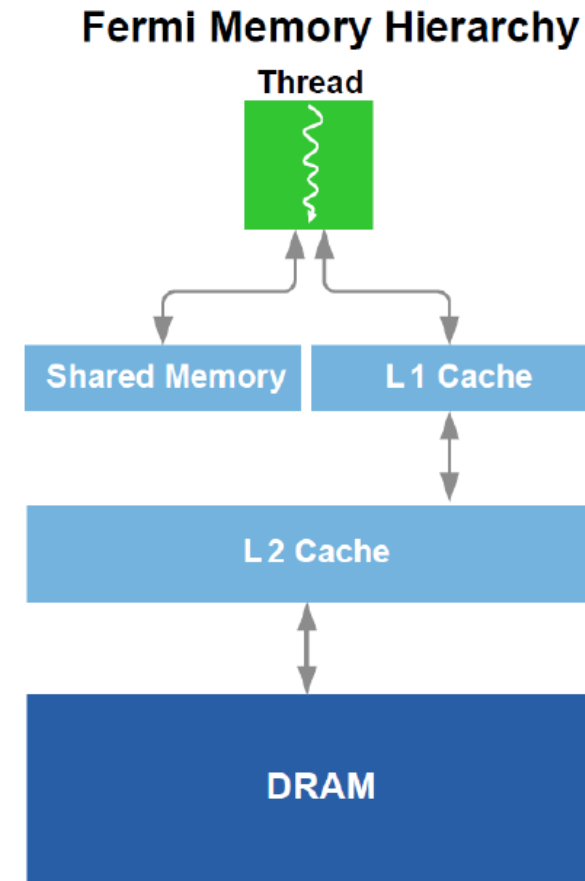
CONFIGURABLE ON-CHIP STORAGE

64 KB per SM

48 KB shared + 16 KB L1
or
16 KB shared + 48 KB L1

Unified L2: 768 KB

Tesla C2050: 14 SMs, 448 cores, 1.03 TFLOPS, 3 GB GDDR5, 144 GB/s



Coalescing groups a warp's memory requests

16

Useful bytes and requested cache sectors can differ by a large factor.

WARP ADDRESS VECTOR



TRANSACTION COUNT FOR 32 THREADS × ONE 4-BYTE WORD

Contiguous

addresses 0, 4, 8, ..., 124

0–31 B

32–63 B

64–95 B

96–127 B

4 sectors = 128 B requested / 128 B useful

Stride ≥ 32 bytes

each lane lands in a different sector



... 32 sectors

32 sectors = 1024 B requested / 128 B useful

These counts describe requested cache sectors. Cache hits and reuse determine how much traffic reaches HBM.

WORKLOAD MILESTONE III

Deep Learning Training

How can we train large models with dense matrix work across GPUs?

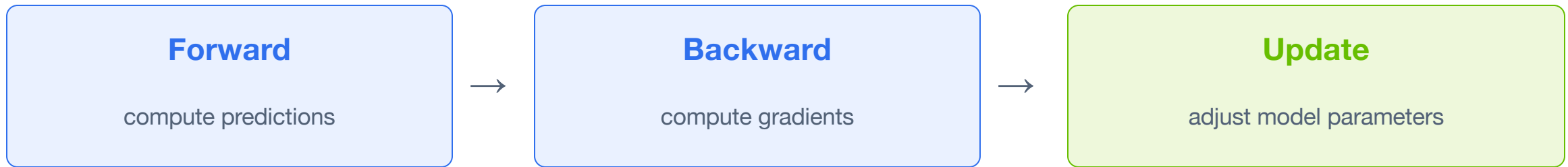
Workload	Repeat forward, backward, and weight-update computations.
Bottleneck	Large matrix operations, training state, and GPU exchanges limit scale.
Architectural response	Use Tensor Cores, high-bandwidth memory, and accelerator links.

Matrix specialization: Volta through Hopper

Training rewards matrix throughput and cooperation

18

Large batches reuse weights; training state and cross-GPU exchanges add memory and link pressure.



Updated weights feed the next training iteration.

Matrix arithmetic

Large matrix operations motivate Tensor Cores and suitable low precision.

Resident state

Weights, activations, gradients, and optimizer state need HBM.

Distributed work

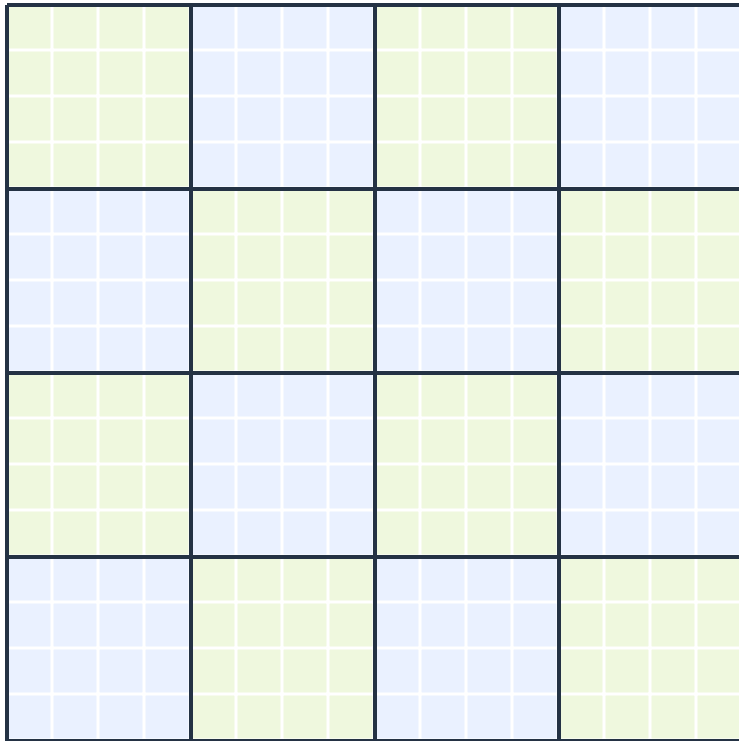
Gradient and partial-result exchanges motivate NVLink and NVSwitch.

Matrix tile size and Tensor Core work

20

Count the arithmetic separately from the hardware cycles that execute it.

OUTPUT TILE D: 16×16 , WITH $K = 16$



$256 \text{ outputs} \times 16 \text{ terms} = 4,096 \text{ FMAs}$

WARP INSTRUCTION

mma / wmma

32 threads collectively supply fragments of A and B and collectively receive fragments of C/D.

WMMA hides fragment layout. PTX documents thread mappings for specific MMA forms.

One tile contains many multiply-adds

CONCRETE HISTORICAL ANCHOR

Volta Tensor Core

$4 \times 4 \times 4 \text{ MMA}$

A 4×4 tile $\times$ a 4×4 tile, accumulated into a 4×4 result.

64 FMA operations / clock

FP16 products, FP32 accumulation

16 output regions $\times$ 4 updates across $K = 64$ small MMA operations. This is a work count, not a cycle count.

Copying data before computation

A consumer must wait for its input, but independent work can continue.

Conventional staging	Global load writes a register; a shared-memory store reads that value.
Direct copy	Dedicated hardware moves the payload from global to shared memory.
Asynchronous execution	Launch the copy, do independent work, then wait before reading the destination.
Two different properties	Direct describes the data path. Asynchronous describes completion timing.

The register saving comes from the direct path, not from asynchrony alone.

Higher bandwidth and direct copies feed matrix work

22

HBM supplies more bytes while asynchronous copies reduce staging overhead.

2016 · PASCAL P100

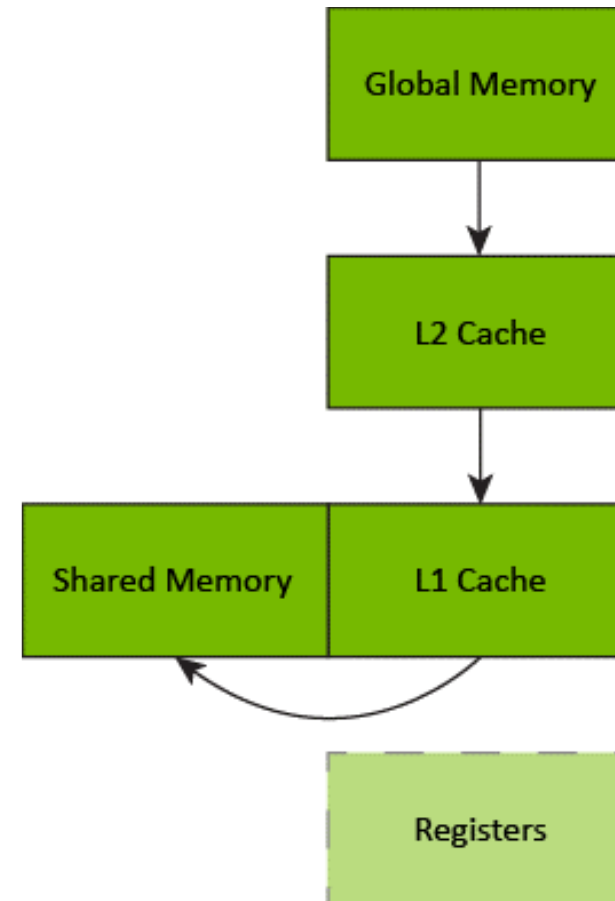
HBM2 moved memory onto the package

16 GB HBM2 at 732 GB/s increased usable bandwidth. NVLink accelerated GPU-to-GPU and GPU-to-CPU transfers. Hardware page faults enabled demand-paged Unified Memory.

2020 · AMPERE A100

Software could stage data without register detours

A100 80 GB SXM: 2.039 TB/s HBM2e and 40 MB L2. Direct global-to-shared copies bypass payload registers. Threads can continue, but must wait before using the copied data.

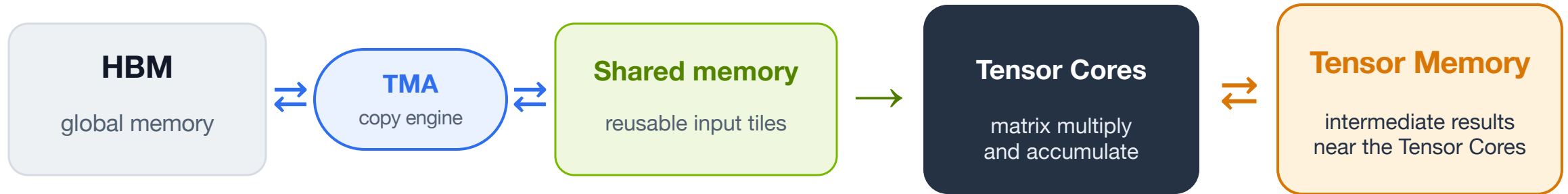


TMA moves tiles and Tensor Memory stores results

23

A transfer engine and a storage resource solve different problems.

INPUT PATH



2022 · HOPPER

TMA first appears

One thread can launch an asynchronous 1D–5D transfer. Hardware handles addresses, strides, and bounds while other threads compute.

2024–2025 · BLACKWELL

Tensor Memory joins the SM

Blackwell Ultra provides 256 KB of TMEM per SM for matrix intermediates. Specialized instructions manage and access it.

2026 · RUBIN

TMA descriptors become easier to reuse

A kernel can override pointer or stride fields, reusing a descriptor for expert tensors with a common layout.

Tensor Core instruction families and operand storage

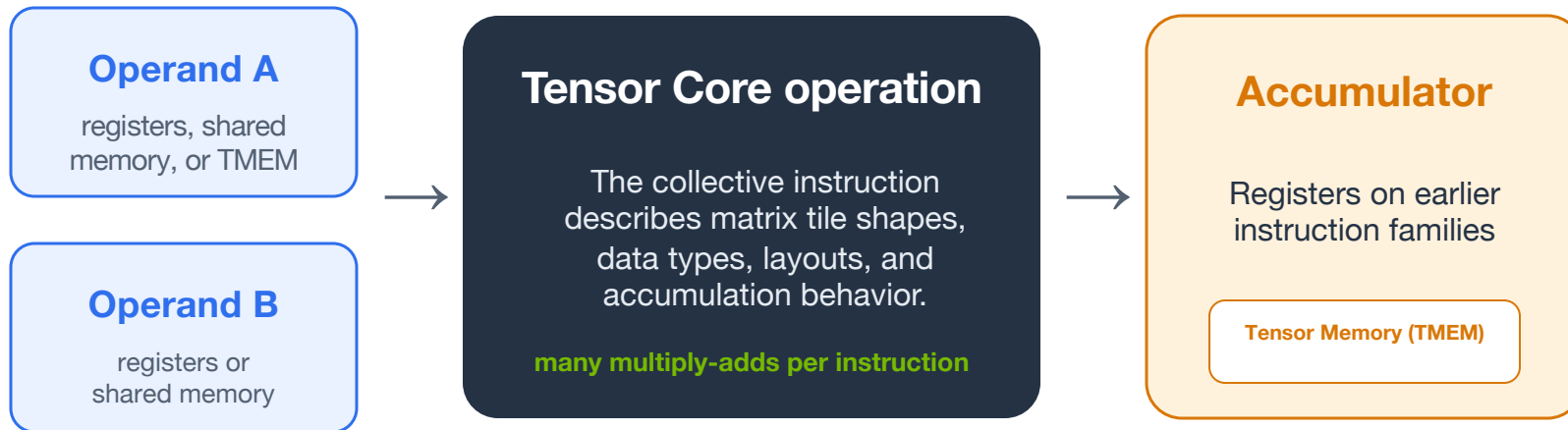
24

The matrix operation stays the same, while participation and operand locations change.

$$D = A \times B + C$$

A and B are input tiles. C is the old accumulator. D is the updated accumulator.

SOURCES VARY BY INSTRUCTION



INSTRUCTION FAMILIES

Volta–Ampere

mma

A warp cooperates. Operands and accumulators use registers.

Hopper

wgmma.mma_async

A 128-thread warpgroup launches asynchronous matrix work. Shared memory can supply operands.

Blackwell

tcgen05.mma

An elected thread launches matrix work. TMEM holds the accumulator. A may use TMEM, B uses shared memory.

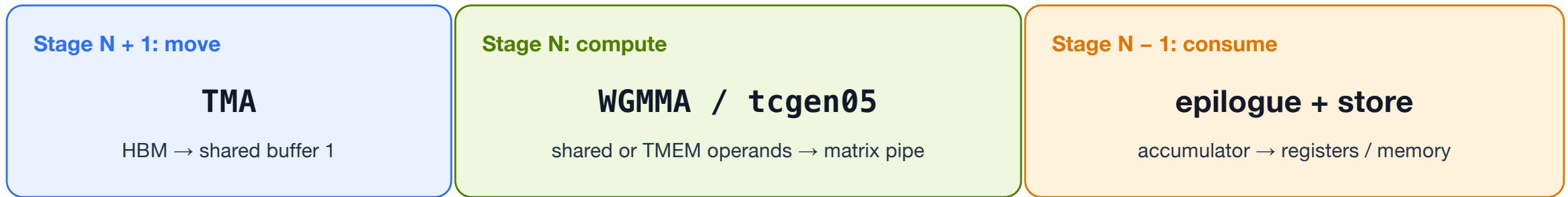
A tile operation coordinates many multiply-adds. One issuing thread does not perform all the arithmetic.

Overlapping tensor copies and matrix computation

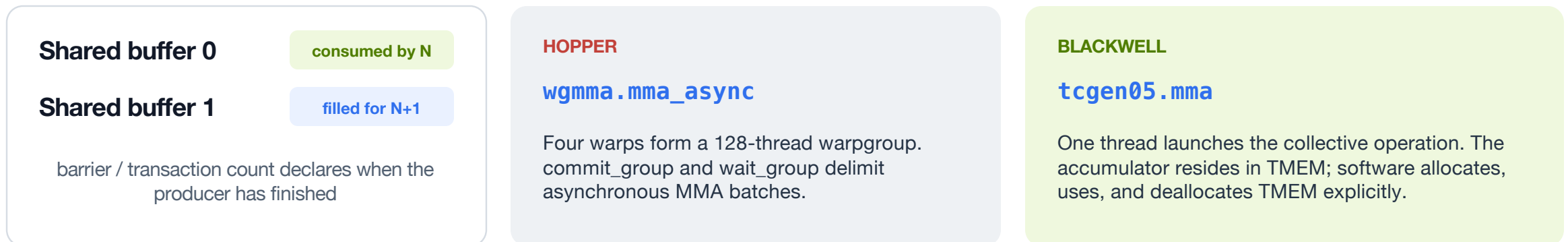
25

Two buffers let the next input tile arrive while the current tile is in use.

STEADY-STATE PIPELINE



BUFFER AND SYNCHRONIZATION STATE



Wait for input arrival before reading. Wait for matrix completion before reusing buffers or consuming results.

Cooperation across SMs and across GPU dies

26

Hopper adds shared-memory cooperation while Blackwell connects two compute dies.

2022 · HOPPER H100 SXM

A new level between shared memory and L2

Thread-block clusters can access distributed shared memory across cooperating SMs. The Tensor Memory Accelerator moves multidimensional tiles while threads compute.

80 GB HBM3

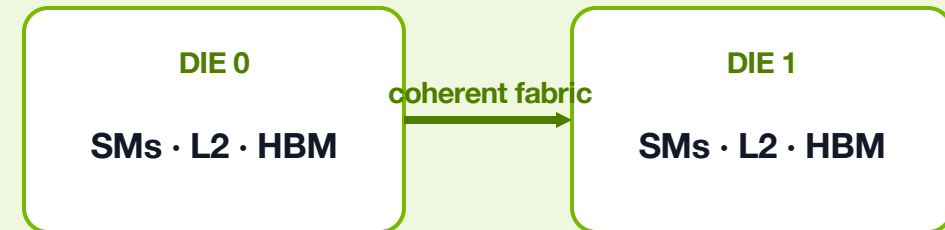
3.35 TB/s HBM bandwidth

50 MB L2 · up to 228 KB shared memory per SM

2025 · BLACKWELL ULTRA B300

Two dies act as one CUDA GPU

Each locality domain has nearby L2 slices and HBM controllers. Coherent L2 caching preserves one address space, but remote traffic crosses the die-to-die fabric.



B300: 288 GB HBM3e and 8 TB/s

160 SMs, 640 Tensor Cores, 256 KB Tensor Memory per SM

WORKLOAD MILESTONE IV

Large Language Model Inference

What limits response time once matrix multiplication becomes fast?

Workload	Process a prompt, then generate dependent tokens one step at a time.
Bottleneck	Small-batch decode can spend more time moving data than multiplying.
Architectural response	Reduce traffic and improve memory, attention, and communication.

Prefill and decode

Inference uses learned weights without updating them.

Prefill	Process the known prompt positions together. Reuse weights across positions.
Decode	Generate the next token, then use it to begin the following step.
Retained state	Keep earlier attention keys and values in the KV cache.
Changing balance	Small-batch decode often offers less weight reuse than prefill.

Batch size and context length change the balance; these are common tendencies.

The attention computation

Compare a query with keys, turn scores into weights, then combine values.

1 Compare

Q contains queries. K contains keys. Row dot products form scores $QK^T / \sqrt{d}$.

2 Normalize

Softmax turns each row of scores into nonnegative weights that sum to one.

3 Combine

Multiply the weights by V, the values, to form the attention output.

d = number of components in each key. Softmax includes exponentials and reductions.

What can delay the next token

Matrix work, other arithmetic, memory access, and communication compete for time.

Which resource or dependency sets the completion time?

1

Matrix computation

Matrix multiplications

2

Memory movement

Weights and KV cache

3

Other arithmetic

Exponentials, reductions,
and scaling

4

Communication

Collectives across
accelerators

Required bytes $\leq$ available bytes. Ideal time $\geq \max(\text{work} / \text{compute rate}, \text{bytes} / \text{bandwidth})$.

What moves? How much? How often? Through what path? Who waits?

Why small-batch decode can be memory limited

Compare the ideal arithmetic time with the time needed to read the operands.

Prefill: known prompt

token 1 token 2 token 3 token 4

Process many positions together.
Often compute-intensive when reuse is high.

Decode: dependent new tokens

new token → next token → next token

Read weights and stored attention keys/values.
Small batches often face memory limits.

Illustrative single-sequence decode step

8B matrix weights × 2 bytes = 16 GB; add 1 GB of old KV reads. All state fits in device memory.

Ideal matrix time

16 GFLOPs / 500 TFLOP/s = 0.032 ms

Ideal read time

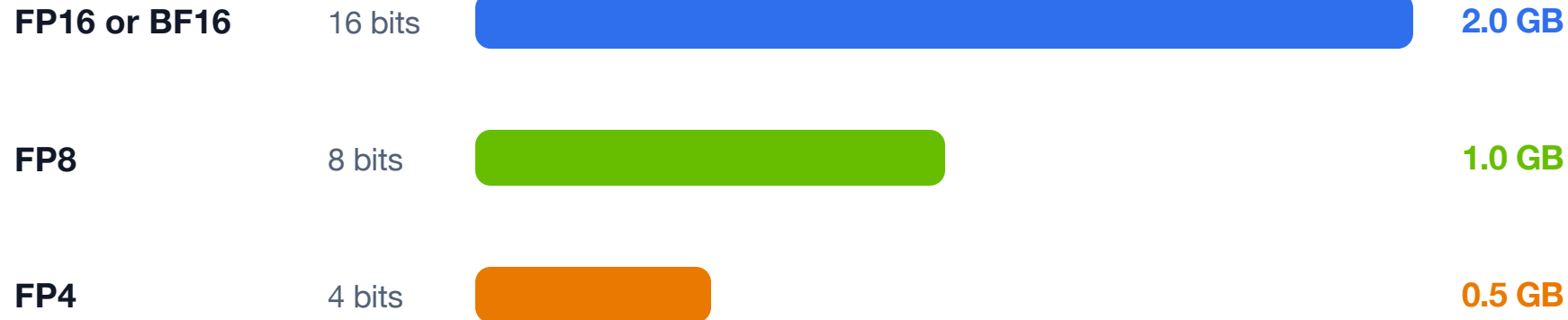
17 GB / 2 TB/s = 8.5 ms

About 266× longer to read. More matrix units alone cannot remove this limit.

Low precision reduces bytes and increases matrix throughput

The same element count can use fewer bytes without reducing the arithmetic count.

STORAGE FOR ONE BILLION VALUES



Compute effect

The same multiply-add count can execute at a higher supported rate.

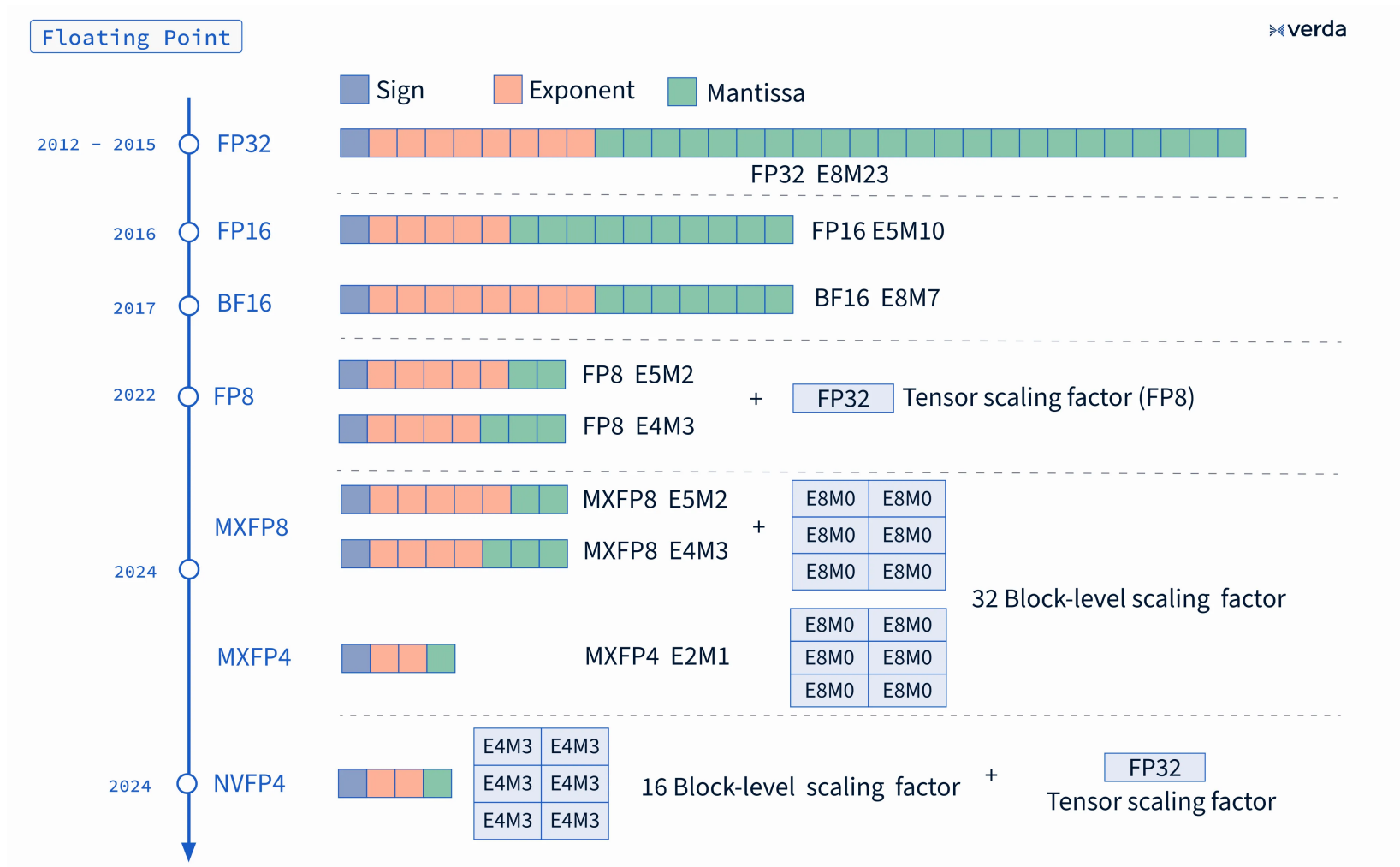
Memory effect

Value-only bytes = element count × bits per value / 8.
Scales and metadata add space.

Value bits only. Decimal GB. Encoding scales, metadata, and padding are additional.

Low precision reduces bytes and increases matrix throughput

The same element count can use fewer bytes without reducing the arithmetic count.



Why softmax matters more after faster matrix multiplication

A tenfold improvement in one phase can produce a much smaller total speedup.

ILLUSTRATIVE TIME: BLUE = GEMM, ORANGE = SOFTMAX



After 10× faster GEMM



Softmax share

10% → 53%

Total speedup = $100 / (9 + 10) = 5.3\times$

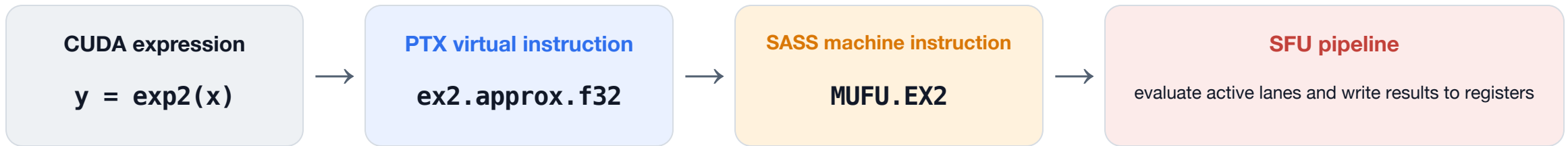
Softmax grows from 10% to 53% of the total. Faster exponentials can reduce part of that cost, while reductions and data movement still matter.

Special function units and exponential throughput

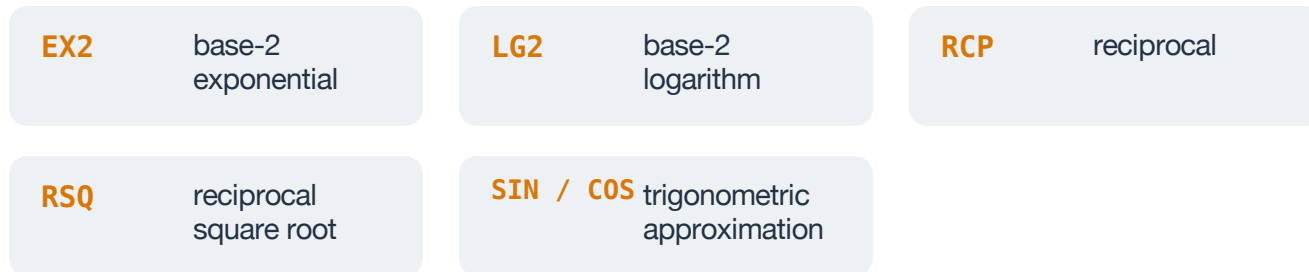
34

Exponentials use a specialized path that can limit softmax performance.

INSTRUCTION PATH



COMMON SFU MACHINE OPERATIONS

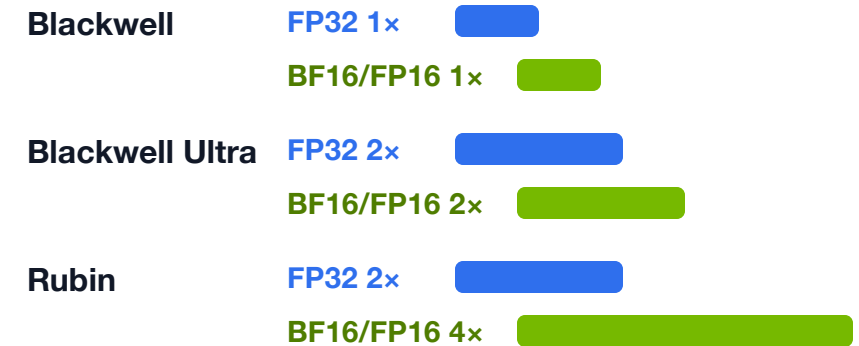


ACCURACY AND SCHEDULING

The hardware instruction supplies an approximate primitive. A compiler or math library may add range reduction and refinement when the requested function needs more accuracy.

EXPONENTIAL THROUGHPUT PER SM

per clock per SM, relative to Blackwell



How a fast exponential uses the SFU

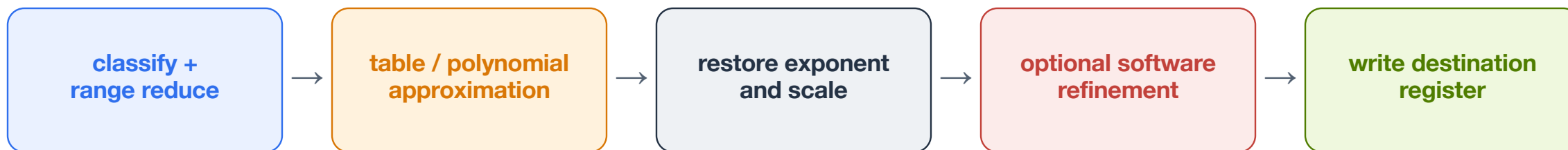
35

A change of base connects $\exp(x)$ to the base-two exponential primitive.

NATURAL EXPONENTIAL EXAMPLE



CONCEPTUAL MATH-FUNCTION PIPELINE



WORKED EXAMPLE

$x = 1$
 $t = 1 \times 1.4427$
 $2^t \approx 2.7183 \approx \exp(1)$

WHAT THE DIAGRAM MEANS

Reduce the input range, approximate, then restore scale.
The boxes explain the mathematics.
The internal NVIDIA circuit is not public.

Matching instructions to execution units

A kernel needs the right mix of arithmetic, matrix, special-function, and memory capacity.

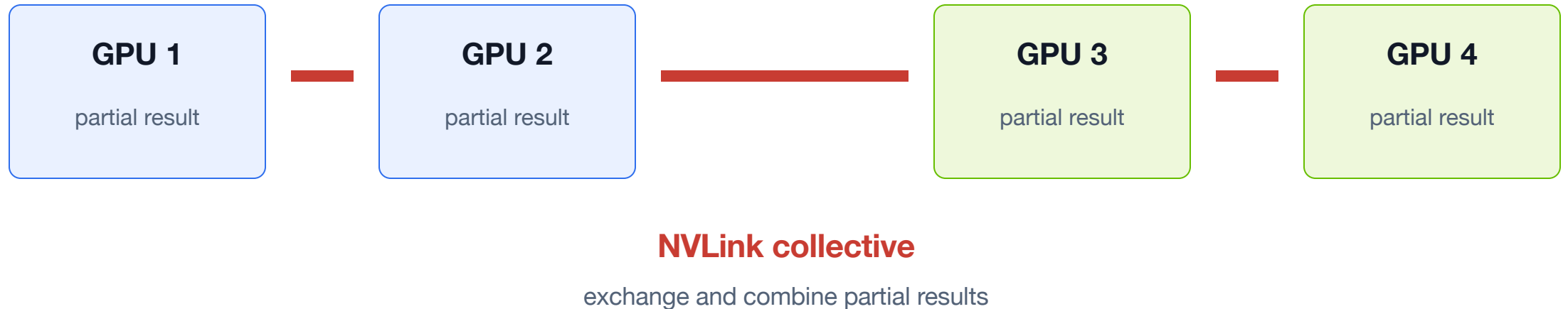
Resource	Work it performs	Typical instructions or operations	Common uses
CUDA cores	General arithmetic across active warp lanes	Floating-point fused multiply-add, integer arithmetic, compare	Elementwise math, indexing, control calculations
Tensor Cores	Matrix multiply-accumulate on small tiles	MMA, WGMMA, and newer Tensor Core instructions	Matrix multiplication, convolution, attention projections
SFUs	Transcendental and special math	Exponential, sine, cosine, reciprocal, reciprocal square root	Softmax exponentials, activations, normalization
Load/store units	Address generation and memory requests	Loads, stores, and atomic operations	Feed registers and shared memory, then write results
Texture units	Cached sampling and filtering	Texture fetch and interpolation	Graphics textures and read-only spatial lookup

NVIDIA uses SFU for Special Function Unit. Warp schedulers coordinate the units but do not perform their arithmetic.

Communication can delay the next layer

GPUs must combine dependent results before the next computation can use them.

ONE DISTRIBUTED LAYER, LOGICAL COMMUNICATION VIEW



If compute time falls faster than communication time, GPUs spend a larger fraction of each layer waiting for data from their peers.

Example: 8 ms compute + 2 ms exchange becomes $4 + 2 = 6$ ms after compute rate doubles.

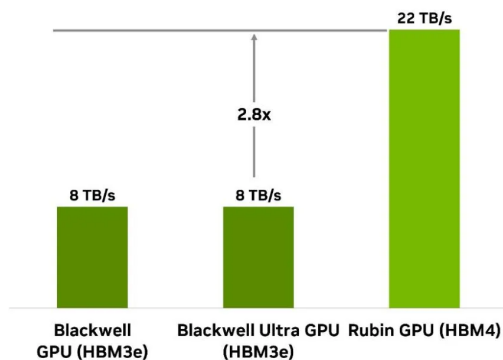
Rubin memory and communication bandwidth

38

Different links serve local memory, neighboring GPUs, and the host.

Rubin Delivers Up To 2.8x More Memory Bandwidth

Blackwell vs Blackwell Ultra vs Rubin



Rubin GPU: 224 SMs and 896 Tensor Cores

ON-PACKAGE MEMORY

Up to 288 GB HBM4
Up to 22 TB/s peak bandwidth

MOVEMENT ENGINES AND CACHES

TMA supports inline descriptor updates. The centralized L2 retains data on chip. These mechanisms reduce movement and setup costs.

SCALE-UP FABRIC

NVLink 6: 3.6 TB/s aggregate per GPU to the switch

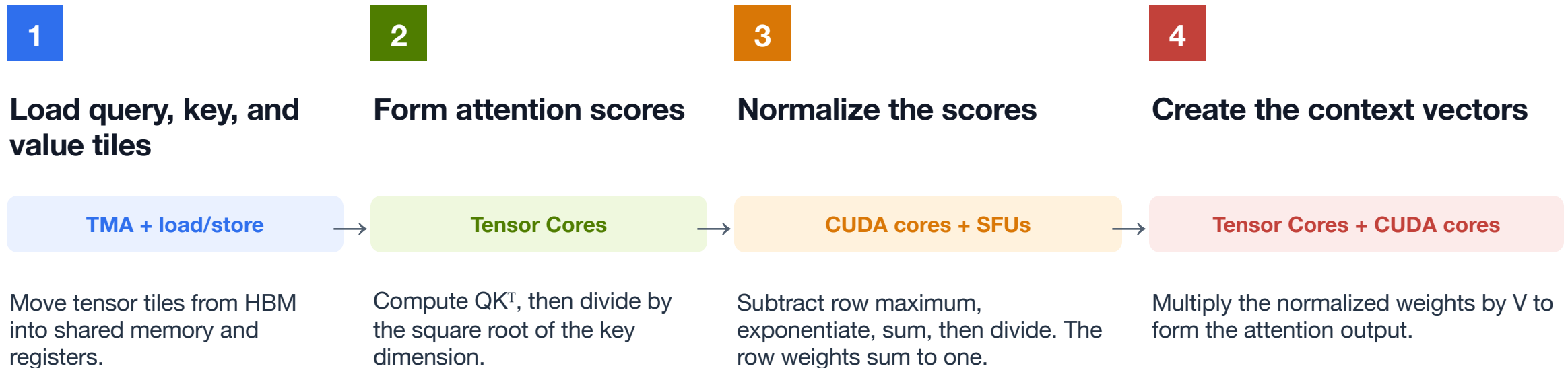
Design lesson

Modern GPU optimization means choosing both a storage level and a movement mechanism for every reused tensor tile.

Attention revisited: mapping work onto GPU hardware

39

Return to the attention algorithm and identify the hardware that serves each stage.



STORAGE USED THROUGHOUT

Registers hold thread-local values. Shared memory stages reusable tiles. Tensor Memory can hold Tensor Core intermediate results on Blackwell-class hardware.

Exact instruction placement depends on the GPU architecture, compiler, and library implementation.

WORKLOAD MILESTONE V

Agentic Computing

What should the machine optimize when the model must act and wait?

Workload

Alternate model execution with tools and observations.

Bottleneck

Persistent state and irregular model–tool dependencies add waiting.

Architectural response

Explore closer integration of compute, memory, and task scheduling.

Emerging workload: model and tool execution

Agentic computing adds tools and persistent state

The complete task alternates between model execution and events outside the GPU.

Heterogeneous execution

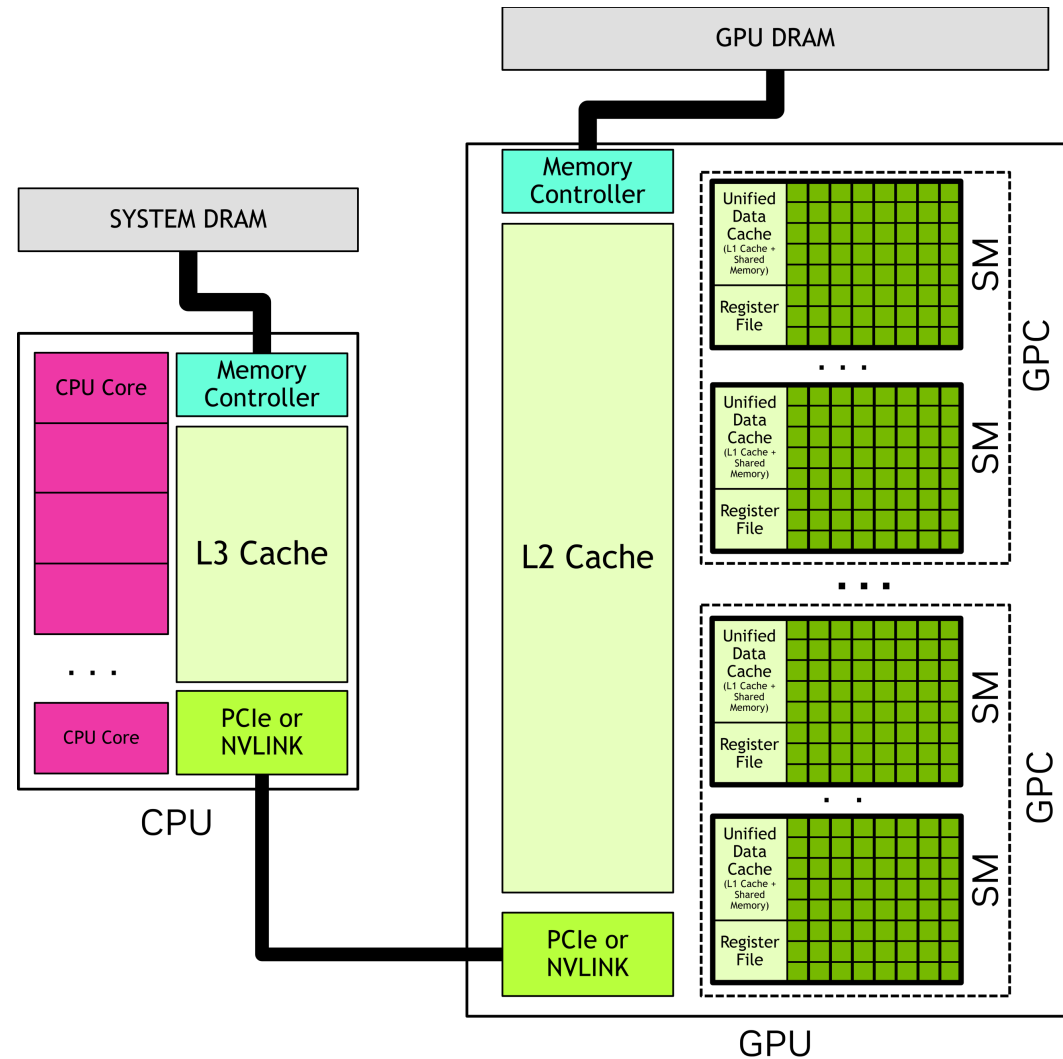
GPU model work alternates with CPU tools, storage, and networks.

Persistent, irregular state

Context and tool results survive pauses, branches, retries, and verification.

Architectural question

Can closer CPU–GPU and memory integration reduce task-level waiting?



Workloads reshape architecture

Workload changes → bottleneck changes → architectural changes

Workload	Pressure on the machine	Architectural response
Graphics	Many similar scene operations	Parallel arithmetic and rendering stages
Scientific computing	General numerical data and cooperation	CUDA threads, local storage, synchronization
Training	Dense matrix work, state, exchanges	Tensor Cores, HBM, accelerator links
Inference	Weight/context traffic and attention	Less traffic, more bandwidth, balanced execution
Agentic computing	Tools, persistent state, irregular waits	Tighter integration? Measure complete tasks.

These workloads coexist. Architectural responses also enable new workloads.

If agentic AI becomes dominant, what should the computer look like?



Reference: GPU evolution in numbers

43

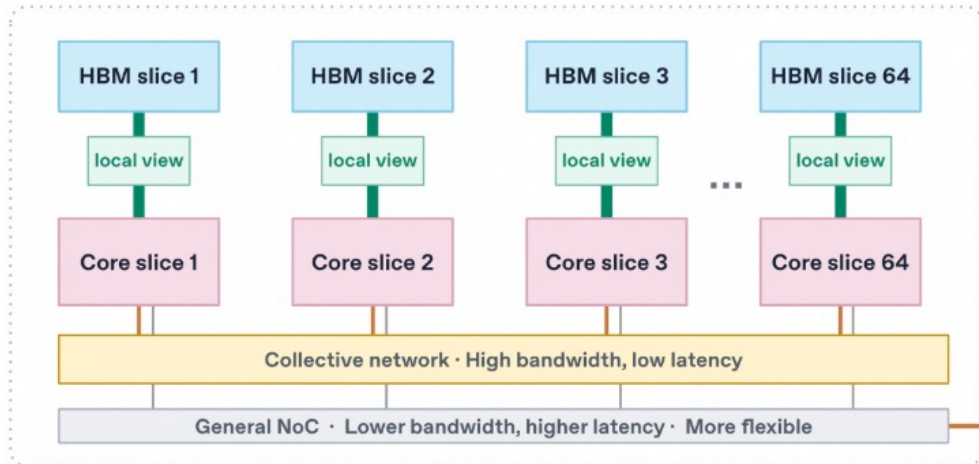
Capacity, bandwidth, and arithmetic throughput describe different limits.

Product	Parallel structure	Peak FP32	Memory	Bandwidth	On-chip storage anchor
GeForce 8800 GTX	128 stream processors	about 0.52 TFLOPS	768 MB GDDR3	86.4 GB/s	16 KB shared memory per SM
Tesla C2050	14 SMs, 448 cores	1.03 TFLOPS	3 GB GDDR5	144 GB/s	64 KB L1/shared per SM, 768 KB L2
Tesla P100 SXM2	56 SMs, 3,584 cores	10.6 TFLOPS	16 GB HBM2	732 GB/s	256 KB registers per SM, 4 MB L2
A100 80 GB SXM	108 SMs, 6,912 cores	19.5 TFLOPS	80 GB HBM2e	2.039 TB/s	192 KB L1/shared per SM, 40 MB L2
H100 SXM	132 SMs, 16,896 cores	67 TFLOPS	80 GB HBM3	3.35 TB/s	up to 228 KB shared per SM, 50 MB L2
Blackwell Ultra B300	160 SMs, 640 Tensor Cores	75 TFLOPS	288 GB HBM3e	8 TB/s	256 KB Tensor Memory per SM
Rubin GPU	224 SMs, 896 Tensor Cores	130 TFLOPS	288 GB HBM4	22 TB/s	centralized L2, capacity not public

B300 and Rubin both list 288 GB. Bandwidth rises from 8 to 22 TB/s, a 2.75× peak-rate increase.

Jalapeño: local memory paths and explicit placement

64 compute slices pair with local HBM regions. Software places data near the computation that uses it. Local tensors, explicit communication, and predictable synchronization.



A shorter local memory path

No shared L2 appears in this diagram. Local HBM access avoids a chip-wide shared-cache path.

Two communication paths

Collective network: common cross-core exchanges.
General NoC: flexible traffic and access to Ethernet.

Per chip: 216 GiB HBM4 at 15.4 TB/s

Matrix peak: 13.4 PFLOP/s for MXFP4 × MXFP4

Trade-off: software must plan data placement and cross-slice communication.

