

# Lecture 6: Hardware and Hardware Algorithms 1

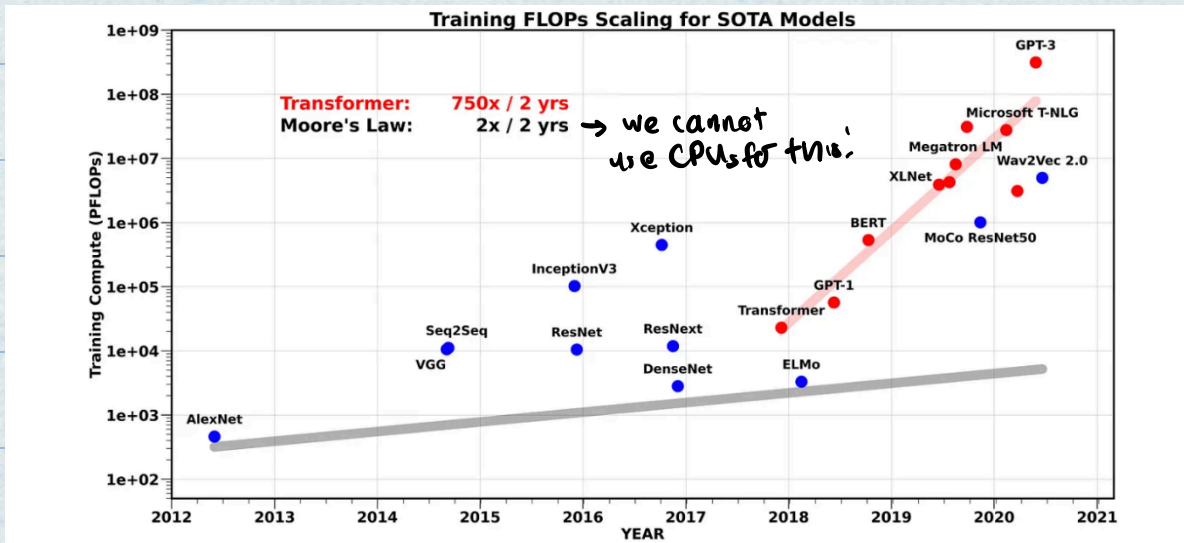
## Agenda:

- CPU vs. GPU overview
- understanding perf.

\* This lecture was entirely taken from Stanford's CS229s

[https://docs.google.com/presentation/d/14hK7SmkUNfSEIRGyptFD2bGO7K9sJOTnwjAVg3vvgg6g/edit#slide=id.g286de50af37\\_0\\_237](https://docs.google.com/presentation/d/14hK7SmkUNfSEIRGyptFD2bGO7K9sJOTnwjAVg3vvgg6g/edit#slide=id.g286de50af37_0_237)

\* Rising Demand for compute that Moore's law is not meeting



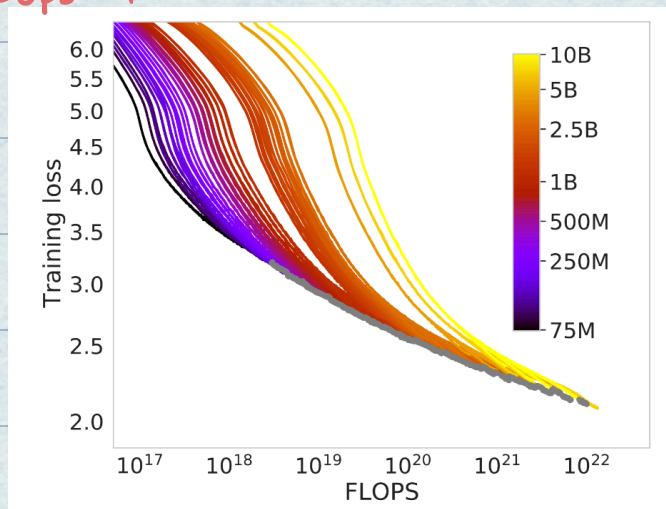
<https://medium.com/riselab/ai-and-memory-wall-2cb4265cb0b8>

\* The more FLOPs the better!

→ given a max # of FLOPs (x-axis) -

what's the best loss we can get? (for

range of model sizes)

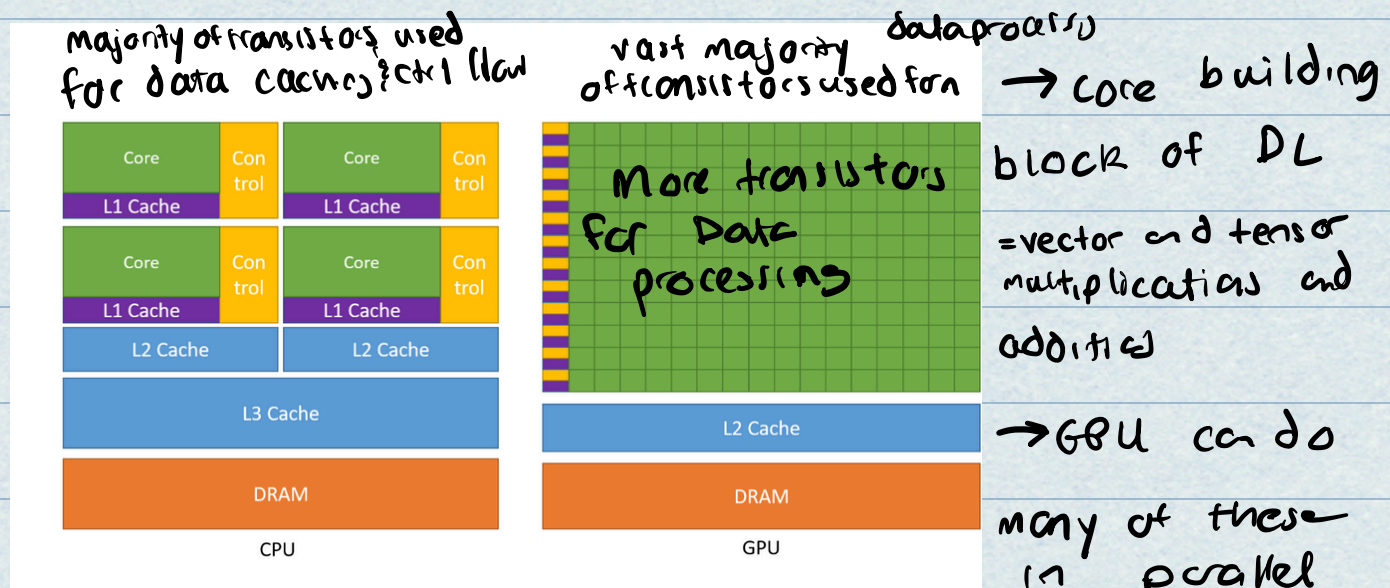


• also: higher model size leads to lower loss



## \* CPU vs. GPU

↳ note there are many dif. new AI accelerators, which we will briefly touch over later



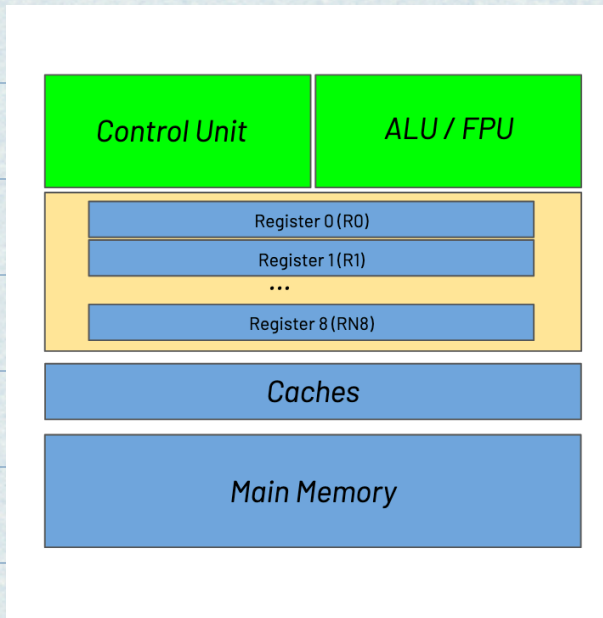
<https://docs.nvidia.com/cuda/cuda-c-programming-guide/index.html#>

- CPUs are designed to execute a seq. of operations, called a thread, w/ minimum latency
- GPUs, in contrast, are designed to excel at executing thousands of threads in parallel
  - generally has slower single-thread performance
  - but achieves significantly high throughput
- CPUs try to use large data caches and complex flow control to hide / avoid long memory access latencies
- GPUs in contrast hide memory access latencies w/ computation.



□ CPU (Central Processing Unit) - Designed for

general-purpose computation



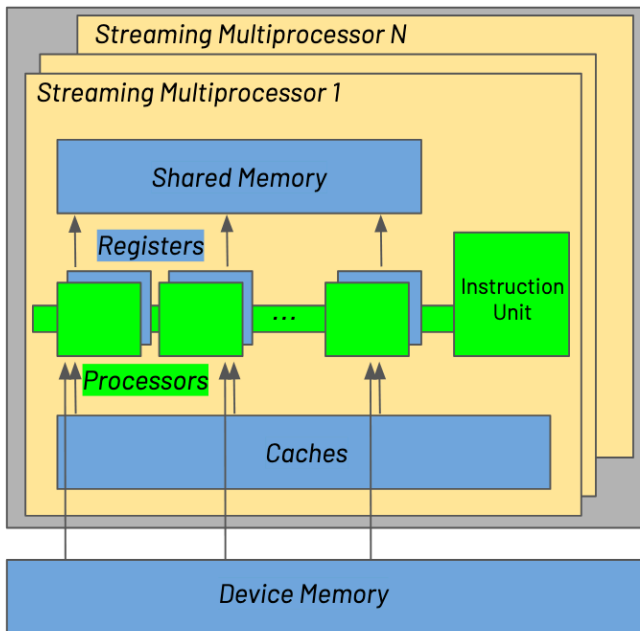
\* Control Unit: fetches next instruction from memory (or cache) and directs arithmetic and floating point unit

\* ALU / FPU: performs bitwise instructions on integer and floating point #s

\* Registers: for next instructions, we can read input unit in and write intermediate values back out

\* Caches / main memory: much larger capacity than registers, used to store instructions and data

□ Graphics Processing Unit (GPU) - for parallelizable programs



\* streaming multiprocessors

each has one instruction unit that controls many processors that run the instructions in parallel

\* registers, shared memory, caches,

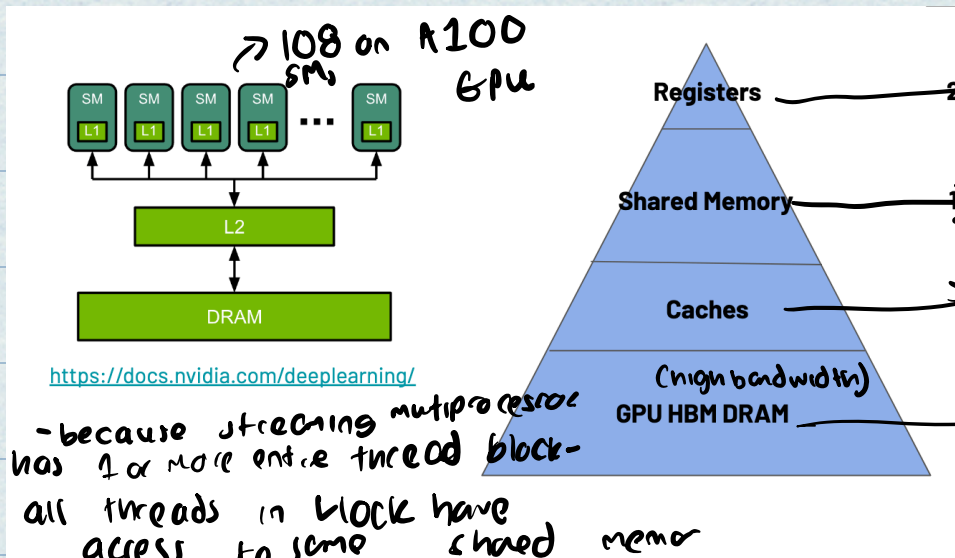
device memory: for reading and writing intermediate data

(more on execution later) - threads are divided into blocks, each SM given 1 or more thread blocks to execute



## \* GPU memory hierarchy:

40 GB A100



<https://docs.nvidia.com/deeplearning/>

- because streaming multiprocessor has 1 or more entire thread blocks - all threads in block have access to same shared memory

Registers → 256 KB/SM

Shared Memory → 192 KB/SM  
↳ bandwidth = 19 TB/s

Caches → 40 MB

(high bandwidth) GPU HBM DRAM → 40 GB

bandwidth = 1.5 - 2 TB/s

\* on-chip SRAM (shared memory)

\* it is key to use

~9x faster than HBM DRAM

on-chip shared memory

but ORDERS of magnitude

effectively when

smaller

designing hardware-aware algorithms

- big difference between CPU and GPU:

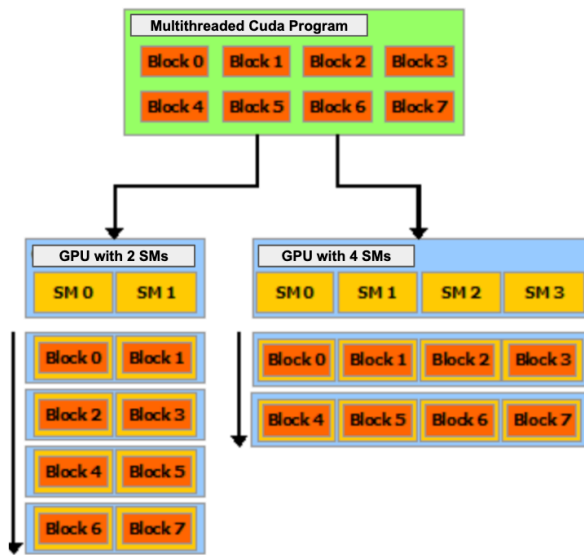
- in GPU programmer has more control of what ends up in each streaming multiprocessor's shared memory

- in contrast, CPU uses caching hierarchy transparently → programmer can only control data layout and order of instructions in program

## \* Programming Nvidia GPUs:



how do we use streaming multiprocessors?



<https://docs.nvidia.com/cuda/cuda-c-programming-guide/>

→ threads in a block can ALSO execute concurrently on a single SM

→ organize

threads into blocks

→ blocks are organized into a GRID

→ blocks of SMs are distributed

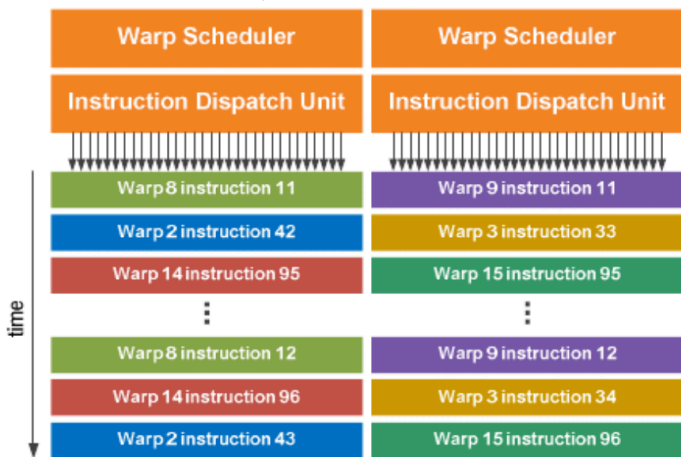
to SMs w/ available execution capacity

→ multiple thread blocks

can execute concurrently on a single SM

\* SMT (Single-instruction, multiple threads)

WARP = 32 parallel threads



- multiprocessor designed to execute hundreds of threads CONCURRENTLY

- each multiprocessor groups <sup>32</sup> threads within a thread block into a WARP

- within a warp: individual threads START at

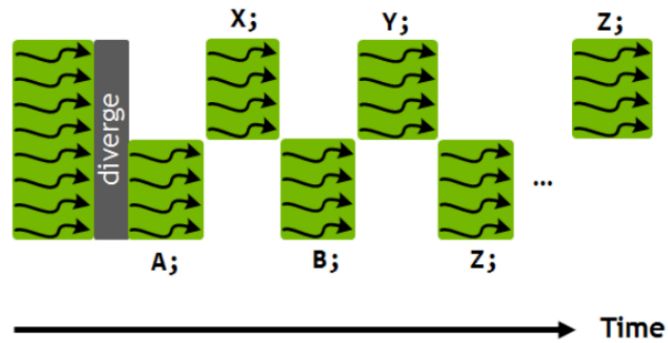
same address, but each thread has its own instruction address counter and register state -  
to can branch independently

↳ but at each time-unit, SM can only execute one instruction at a time, so



optimal for there to be no branches:

```
if (threadIdx.x < 4) {
    A;
    B;
} else {
    X;
    Y;
}
Z;
```



## How an if-else is executed in the Volta implementation of SIMT.

<https://images.nvidia.com/content/volta-architecture/pdf/volta-architecture-whitepaper.pdf>

(Figure 22)

→ Within an SM, all execution context (program counters, registers, etc) are maintained on CTRP until warp is done executing - so "switching" btwn dif. warps on same SM has no cost

↳ at each timestep wop scheduler will try to select a wop that has next thread ready for run

Perf

## \* Part 2 of Lecture: Thinking About

- in alg. class, we think about

algorithms in terms of asymptotics:

↳ Strassen's algorithm for matrix

multiply takes  $O(N^{2.8074})$



↳ naive matrix multiply  $O(N^3)$  for  $N \times N$  matrices

- but big  $O$  notation ignores:

- constants that can make a large diff. for different problem sizes

- implementation specific details!  
(these matter A LOT in real life)

- an algorithm could be ASYMPTOTICALLY BETTER but have worse perf. if not suited for specific hardware

\*concrete example: Algorithmic scaling  $\neq$  realworld perf.

→ 1+ELU\* is a fancy new attn algorithm w/ better asymptotic scaling than normal (softmax) attn.

Method:

softmax

1+ELU

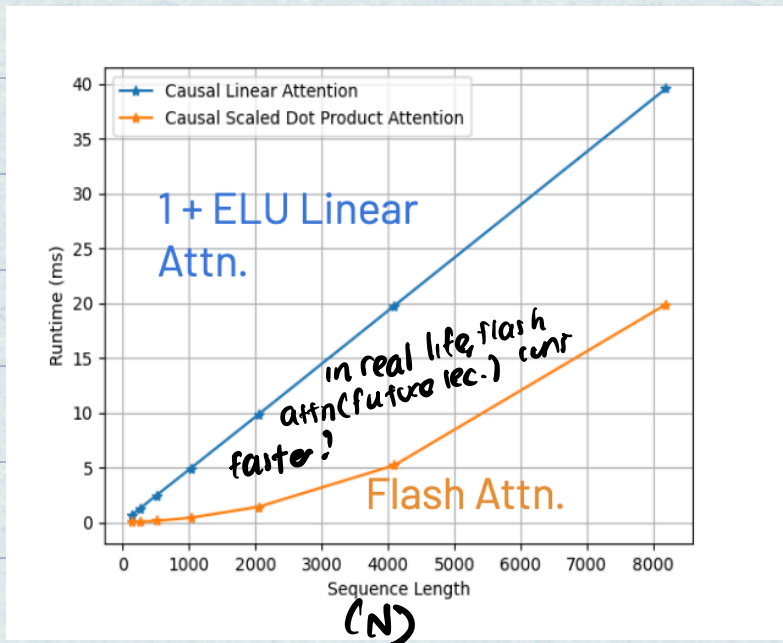
$O(N^2d)$

$O(nd^2)$

↓  
dimension of Q/K/V matrices

input seq. length

complexity



Taken from cs299s lecture, original references:

\* Transformers are RNNs: Fast Autoregressive Transformers with Linear Attention. Katharopoulos et al., ICML 2020. <https://linear-transformers.com/>

\*\*Scaling without linear attention cuda kernel. Custom cuda kernel makes linear attention go faster.

Credit: Michael Zhang <https://michaelzhang.>

\* How do we THINK about performance ???



- alg. could be good asymptotically but not suited for specific hardware

## \* HARDWARE UTILIZATION (consider):

- ↳ - theoretical peak perf.
- memory bandwidth of HW
- try to find FRACTION of peak theoretical perf. attained by an alg
- use this info to decide how to improve perf.

## \* Theoretical Peak Performance

- in any HW (CPU or GPU), we:
  - move items from memory to a processing unit
  - perform some computation
  - move items back to memory
- to get peak perf, we need the following HW properties:

\* memory bandwidth  $\rightarrow$  max # of

items that can be moved to / from memory to / from processor in each second

\* compute bandwidth: max # of ops

that can be done in processor per second



## \* Example peak perf. for NVIDIA A100:

|                        | A100 80GB PCIe           | A100 80GB SXM |
|------------------------|--------------------------|---------------|
| FP64                   | 9.7 TFLOPS               |               |
| FP64 Tensor Core       | 19.5 TFLOPS              |               |
| FP32                   | 19.5 TFLOPS              |               |
| Tensor Float 32 (TF32) | 156 TFLOPS   312 TFLOPS* |               |
| BFLOAT16 Tensor Core   | 312 TFLOPS   624 TFLOPS* |               |
| FP16 Tensor Core       | 312 TFLOPS   624 TFLOPS* |               |
| INT8 Tensor Core       | 624 TOPS   1248 TOPS*    |               |
| GPU Memory             | 80GB HBM2e               | 80GB HBM2e    |
| GPU Memory Bandwidth   | 1,935 GB/s               | 2,039 GB/s    |

→ many larger

flops vs  
memory BW

<https://developer.nvidia.com/blog/nvidia-ampere-architecture-in-depth/>

## \* compute vs. memory bound:

- $T_{mem}$  = time spent accessing memory
- $T_{math}$  = time spent doing math (on compute)
- when memory I/O is overlapped, total execution time is:

$$\max(T_{mem}, T_{math})$$

\* compute bound: when  $T_{math} > T_{mem}$

\* memory bound: when  $T_{mem} > T_{math}$

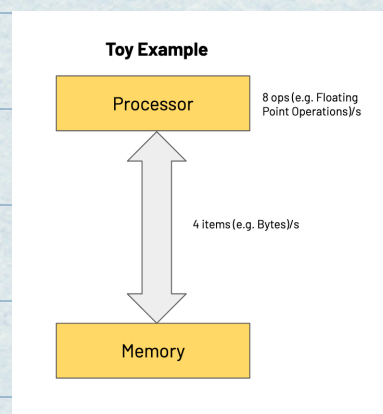
## \* Peak Bandwidth vs. Peak Compute mismatch:

Here:

compute BW = 8 ops/s

mem BW = 4 item/s

- suppose we implement the





following als:

1. Load 8 items from memory into processor.
2. run 1 op on each item
3. write all 8 outputs back

- if compute overlaps w/ I/O, how long does this take to run?

$$\text{Total time} = 2s (\text{load}) + \overset{\text{overlap}}{2s (\text{compute} + \text{stall})} + 2s (\text{write})$$

→ second 0-1: first 4 items arrive

→ 1-2: 4 items computed on, 2nd 4 arrive

→ 2-3: first 4 written back, 2nd 4 computed

→ 3-4: 2nd 4 written back

- for sufficiently large <sup># of</sup> items:

- 4 items arrive per 1s to processor
- processor could do 8 ops, but does 4

↳ 50% compute utilization but 100% mem BW utilization

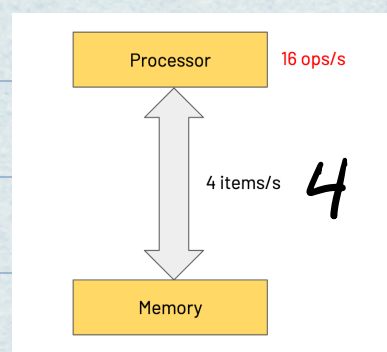
\* Example 2: Increase compute speed to 16 ops/s

→ total time = unchanged!

overlap

$$2s (\text{load}) + \overset{\text{overlap}}{2s (\text{compute} + \text{stall})} + 2s (\text{write})$$

↳ 25% compute utilization



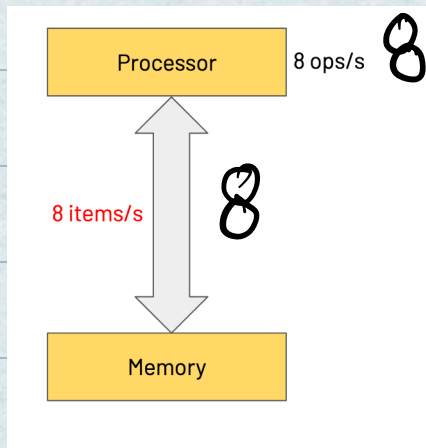
16



and 100 % mem BW utilization,

\* increasing compute speed does NOT improve perf

\* Example 2: Increase mem BW to 8



Total time = 3 overlap  
1s (load) + 1s (compute) + 1s (load)  
3 seconds!

→ 0-1: 8 items arrive

→ 1-2: 8 items computed

→ 2-3: 8 items written

↳ 100 % compute and BW utilization

\* example 4: original HW, but new als

1. Load 4 items from mem

2. compute 4 ops on EACH item

3. write 4 outputs back

total time = 4

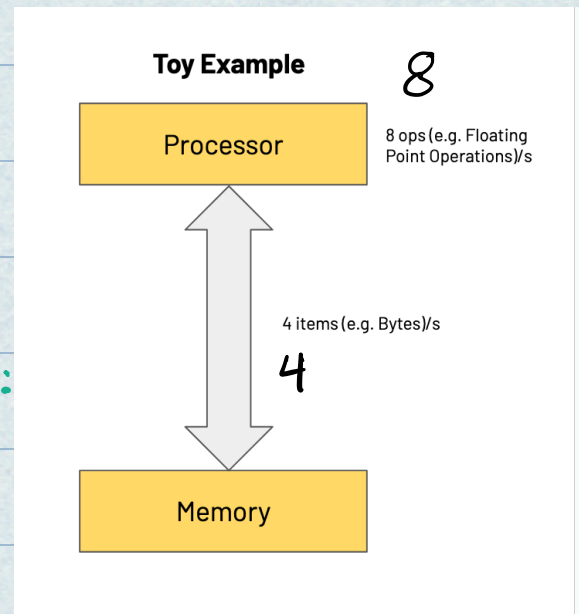
1s (load) + 2s (compute) + 1s (write):

• 4 items arrive at proc. in steady state PER second

• processor needs to do 16

ops per second, but can only do 8

→ 100 % compute AND BW utilization





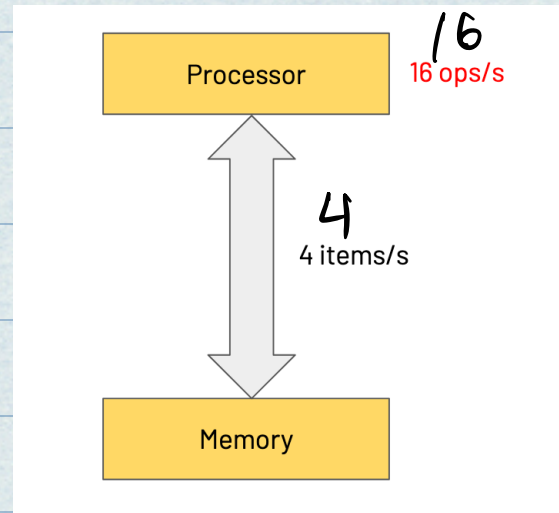
\* example 5: a lg2, but increase compute speed to 16

total time = 3

1s (load) + 1s (compute) +  
1s (write)

→ 4 items arrive per second  
→ processor needs to do  
4 ops/item / second, it can!

→ 100% compute and BW  
utilization



\* Compute Bound vs. Bandwidth Bound

- recall: compute bound when  $T_{\text{math}} > T_{\text{mem}}$
- mem. bound when  $T_{\text{mem}} > T_{\text{math}}$

$$T_{\text{math}} = \frac{\# \text{ ops}}{\text{compute BW}}$$

$$T_{\text{mem}} = \frac{\# \text{ bytes accessed}}{\text{mem BW}}$$

(alternately - think abt ops / byte):

• compute bound:

$$\frac{\# \text{ ops}}{\text{bytes accessed}} > \frac{\text{compute BW}}{\text{mem BW}}$$

• mem bound:

$$\frac{\# \text{ ops}}{\text{bytes accessed}} < \frac{\text{compute BW}}{\text{mem BW}}$$



## \* Arithmetic (Operational) intensity:

→ how many ops do we perform for each byte moved?

$$\rightarrow \text{Arithmetic Intensity} = \frac{\text{Total \# ops}}{\text{\# Bytes written/read to/from mem.}}$$

## \* Lets re-consider A100:

• 1935 GB/s mem-

bw

• 312 TFLOPS for

FP16

• peak flop (or ratio

of compute to mem.

$$\text{bw}) = \frac{312 \text{ FLOPS}}{1935 \text{ GB/s}} = 161 \frac{\text{FLOPS}}{\text{byte}}$$

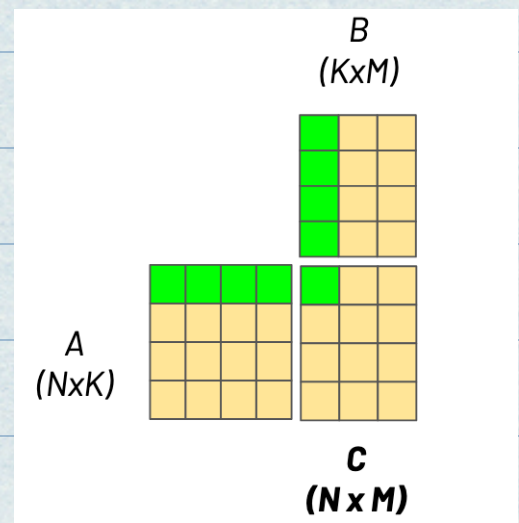
|                        | A100 80GB PCIe           | A100 80GB SXM |
|------------------------|--------------------------|---------------|
| FP64                   | 9.7 TFLOPS               |               |
| FP64 Tensor Core       | 19.5 TFLOPS              |               |
| FP32                   | 19.5 TFLOPS              |               |
| Tensor Float 32 (TF32) | 156 TFLOPS   312 TFLOPS* |               |
| BFLOAT16 Tensor Core   | 312 TFLOPS   624 TFLOPS* |               |
| FP16 Tensor Core       | 312 TFLOPS   624 TFLOPS* |               |
| INT8 Tensor Core       | 624 TOPS   1248 TOPS*    |               |
| GPU Memory             | 80GB HBM2e               | 80GB HBM2e    |
| GPU Memory Bandwidth   | 1,935 GB/s               | 2,039 GB/s    |

\* We can also compute arithmetic intensity of common algorithms:

\* Lets consider matrix multiply



→ recall: output  $C$  has  $N \times M$  entries, each entry comes from dot product of vectors  $(A_i, B_j)$  of dimension  $K$ :



→ each dot product is  $K$  multiplies and  $K$  adds, totaling  $2K$  flops:

$$C_{ij} = A_{i1} \cdot B_{1j} + A_{i2} \cdot B_{2j} \dots A_{iK} \cdot B_{Kj}$$

$$\text{total flop} = 2MNK$$

\* what abt memory?

→ I/O components:

1) read  $A$ :  $N \times K$

2) read  $B$ :  $K \times M$

3) write  $C$ :  $N \times M$

→ for fpl6, each val = 2 bytes

→ total bytes accessed:  
 $2(NK + Km + Nm)$

\* arithmetic intensity of math mul is:

$$\frac{2MNK}{2(NK + Km + Nm)}$$

\* consider 2 examples:

$$\begin{aligned} 1) \quad M = K = 8192 \quad N = 128 &\rightarrow \frac{2(8192 \cdot 128 \cdot 8192)}{2(8192 \cdot 128 + 8192^2 + 128 \cdot 8192)} = 124.1 \end{aligned}$$

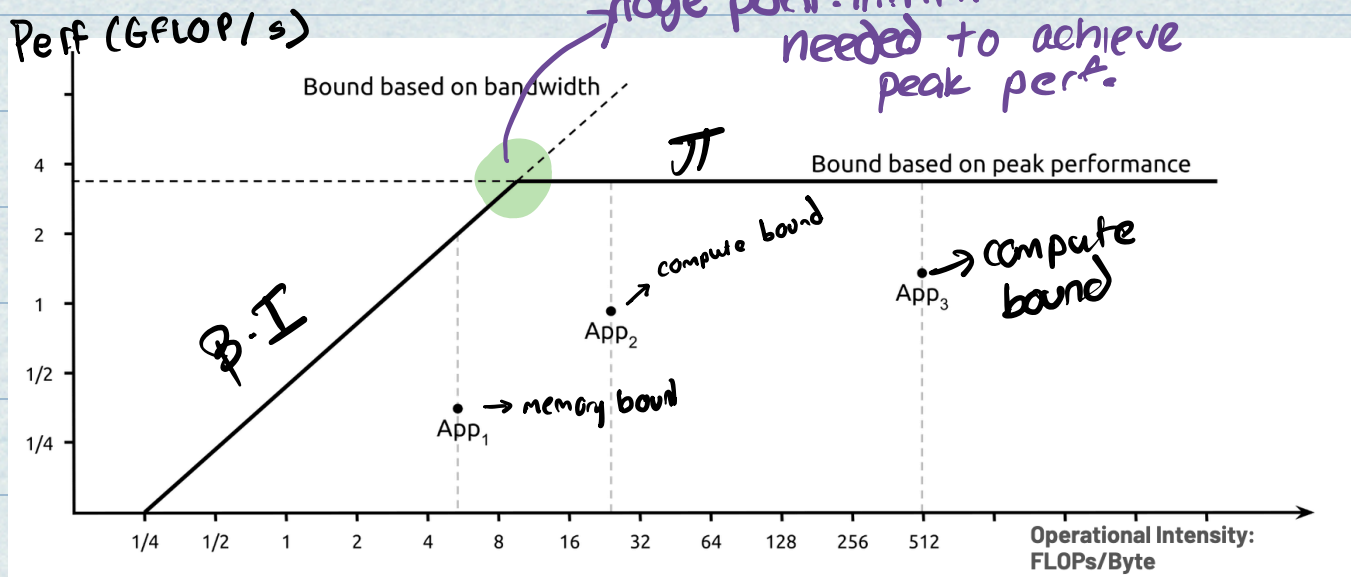
for A100:  $124.1 < 161 \rightarrow$  memory bound

$$2) \quad M = K = N = 8192 \rightarrow \text{arithmetic intensity} = 2730.6$$



2730 .6 > 161 → compute bound!

## \* ROOFLINE MODEL



[https://en.wikipedia.org/wiki/Roofline\\_model](https://en.wikipedia.org/wiki/Roofline_model)

· Idea: plot perf. (GFLOP/s) as function of arithmetic / operational intensity (FLOPs/Byte)

$P = \min \left\{ \begin{array}{l} \pi = \text{peak perf. of this HW} \left( \frac{\text{peak compute BW}}{\text{peak mem BW}} \right) \\ \beta \cdot I = \text{product of mem. BW and} \\ \text{algorithm's arithmetic intensity} \end{array} \right.$

→ next time:

- some simple steps for improving perf.
- analyzing efficiency of transformers