

GP-GPU and High Performances Computing

Cours 1
Introduction

What day are you born ? (keep the day in your head)

Add the digits

If number is greater than 10, keep the last digit

If number is less than 10, keep the digit

Say the number aloud

Find the smallest number in your row - in the room.

TOP 500

Rank	System	Cores	Rmax (PFlop/s)	Rpeak (PFlop/s)	Power (kW)
1	El Capitan - HPE Cray EX255a, AMD 4th Gen EPYC 24C 1.8GHz, AMD Instinct MI300A, Slingshot-11, TOSS, HPE DOE/NNSA/LLNL United States	11,039,616	1,742.00	2,746.38	29,581
2	Frontier - HPE Cray EX235a, AMD Optimized 3rd Generation EPYC 64C 2GHz, AMD Instinct MI250X, Slingshot-11, HPE Cray OS, HPE DOE/SC/Oak Ridge National Laboratory United States	9,066,176	1,353.00	2,055.72	24,607
3	Aurora - HPE Cray EX - Intel Exascale Compute Blade, Xeon CPU Max 9470 52C 2.4GHz, Intel Data Center GPU Max, Slingshot-11, Intel DOE/SC/Argonne National Laboratory United States	9,264,128	1,012.00	1,980.01	38,698
4	JUPITER Booster - BullSequana XH3000, GH Superchip 72C 3GHz, NVIDIA GH200 Superchip, Quad-Rail NVIDIA InfiniBand NDR200, RedHat Enterprise Linux, EVIDEN EuroHPC/FZJ Germany	4,801,344	793.40	930.00	13,088
5	Eagle - Microsoft NDv5, Xeon Platinum 8480C 48C 2GHz, NVIDIA H100, NVIDIA Infiniband NDR, Microsoft Azure Microsoft Azure United States	2,073,600	561.20	846.84	

HPCG 500

Rank	TOP500 Rank	System	Cores	Rmax (PFlop/s)	HPCG (TFlop/s)
1	1	El Capitan - HPE Cray EX255a, AMD 4th Gen EPYC 24C 1.8GHz, AMD Instinct MI300A, Slingshot-11, TOSS, HPE DOE/NNSA/LLNL United States	11,039,616	1,742.00	17406.90
2	7	Supercomputer Fugaku - Supercomputer Fugaku, A64FX 48C 2.2GHz, Tofu interconnect D, Fujitsu RIKEN Center for Computational Science Japan	7,630,848	442.01	16004.50
3	2	Frontier - HPE Cray EX235a, AMD Optimized 3rd Generation EPYC 64C 2GHz, AMD Instinct MI250X, Slingshot-11, HPE Cray OS, HPE DOE/SC/Oak Ridge National Laboratory United States	9,066,176	1,353.00	14054.00
4	3	Aurora - HPE Cray EX - Intel Exascale Compute Blade, Xeon CPU Max 9470 52C 2.4GHz, Intel Data Center GPU Max, Slingshot-11, Intel DOE/SC/Argonne National Laboratory United States	9,264,128	1,012.00	5612.60
5	9	LUMI - HPE Cray EX235a, AMD Optimized 3rd Generation EPYC 64C 2GHz, AMD Instinct MI250X, Slingshot-11, HPE EuroHPC/CSC Finland	2,752,704	379.70	4586.95

Green 500

Rank	System	Cores	Rmax (PFlop/s)	Power (kW)	Energy Efficiency (GFlops/watts)
259	JEDI - BullSequana XH3000, Grace Hopper Superchip 72C 3GHz, NVIDIA GH200 Superchip, Quad-Rail NVIDIA InfiniBand NDR200, ParTec/EVIDEN EuroHPC/FZJ Germany	19,584	4.50	67	72.733
148	ROMEO-2025 - BullSequana XH3000, Grace Hopper Superchip 72C 3GHz, NVIDIA GH200 Superchip, Quad-Rail NVIDIA InfiniBand NDR200, Red Hat Enterprise Linux, EVIDEN ROMEO HPC Center - Champagne-Ardenne France	47,328	9.86	160	70.912
484	Adastra 2 - HPE Cray EX255a, AMD 4th Gen EPYC 24C 1.8GHz, AMD Instinct MI300A, Slingshot-11, RHEL, HPE Grand Equipement National de Calcul Intensif - Centre Informatique National de l'Enseignement Suprieur (GENCI-CINES) France	16,128	2.53	37	69.098
183	Isambard-AI phase 1 - HPE Cray EX254n, NVIDIA Grace 72C 3.1GHz, NVIDIA GH200 Superchip, Slingshot-11, HPE University of Bristol United Kingdom	34,272	7.42	117	68.835

General information

General course information

Contact - Christophe Picard

- christophe.picard@univ-grenoble-alpes.fr
- Offices :
 - 174 – IMAG Building – 1st floor
 - C102 - Ensimag Building - 1st floor
- Email headers: [CGPU]
- Course notes available on my website

Class organization

- Lectures and labs
- One written exam **and** a project

About the project

- It is **your** project. You **choose** the subject.
- It can be a group project (2 or 3).
- It should be related to your major (MMIS, MSIAM, MoSiG).
- You may reused project from other courses, but you should propose an improvement.
- It must involve some programming (Python, C, C++, Julia, Fortran).
- It must involve some design. You are **not allowed** to just use a library for parallelism.
- You must submit your proposal for project by **October 13th 2025**.

Grading the project

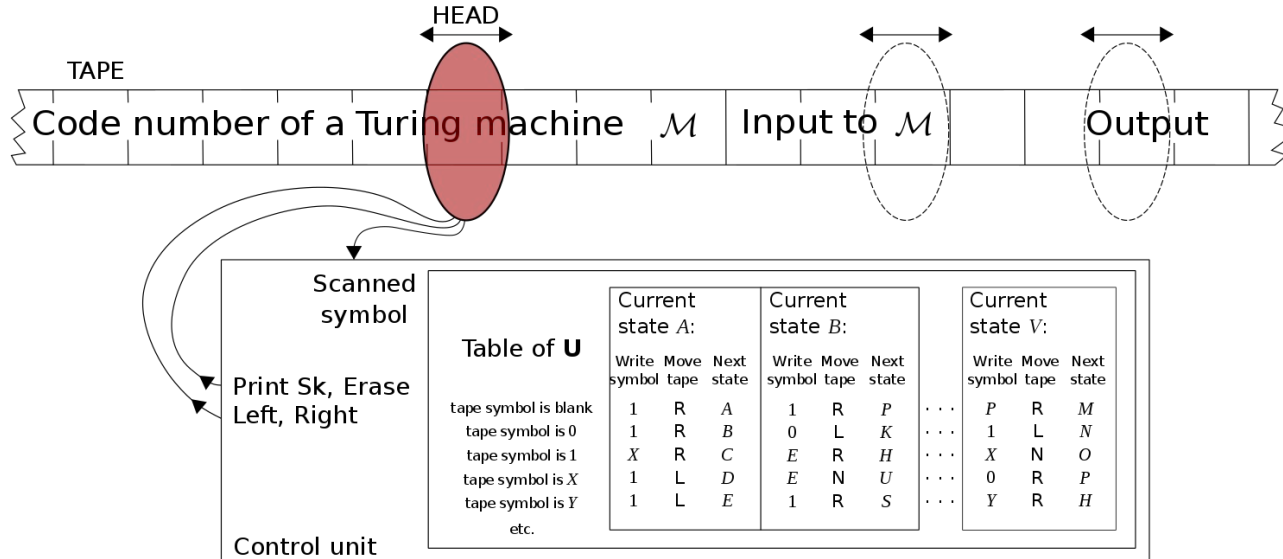
The grade will depend on :

- the quality of the code.
- the quality of the report.
- the design of the parallelism.
- the performances.
- the number of students in the project.
- the difficulty of the subject.

Introduction

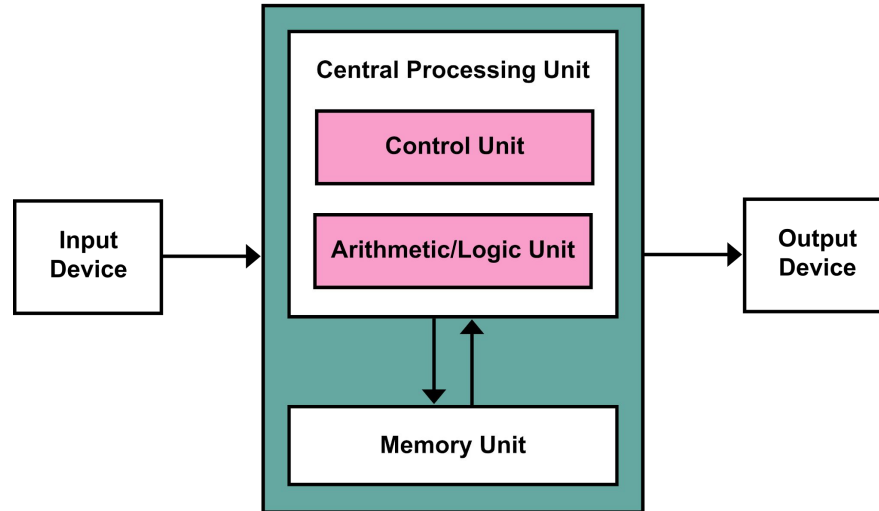
What's a computer look like?

A universal Turing Machine



What's a computer look like?

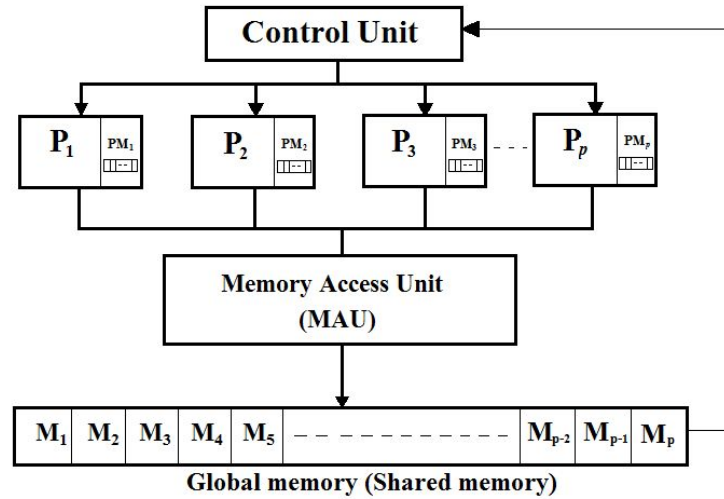
Von Neumann Architecture



Properties of “computers”

- Sequential processing
 - ◆ Control or logical flow
- Algorithm costs measured in this model
 - ◆ Big-O notation counts number of sequential steps/bits for storage
- This is the basis for the CS curriculum
 - ◆ And it's just wrong
 - ◆ Computers are not sequential and performance is more nuanced than counting the number of steps
- We look at computers as parallel entities
 - ◆ Do many **tasks concurrently**
 - ◆ **Tasks interfere** with each other
 - ◆ More accurately reflects hardware and bottlenecks
- What about parallel computation models?
 - ◆ Exist but **not useful**, because reality collides with the abstraction

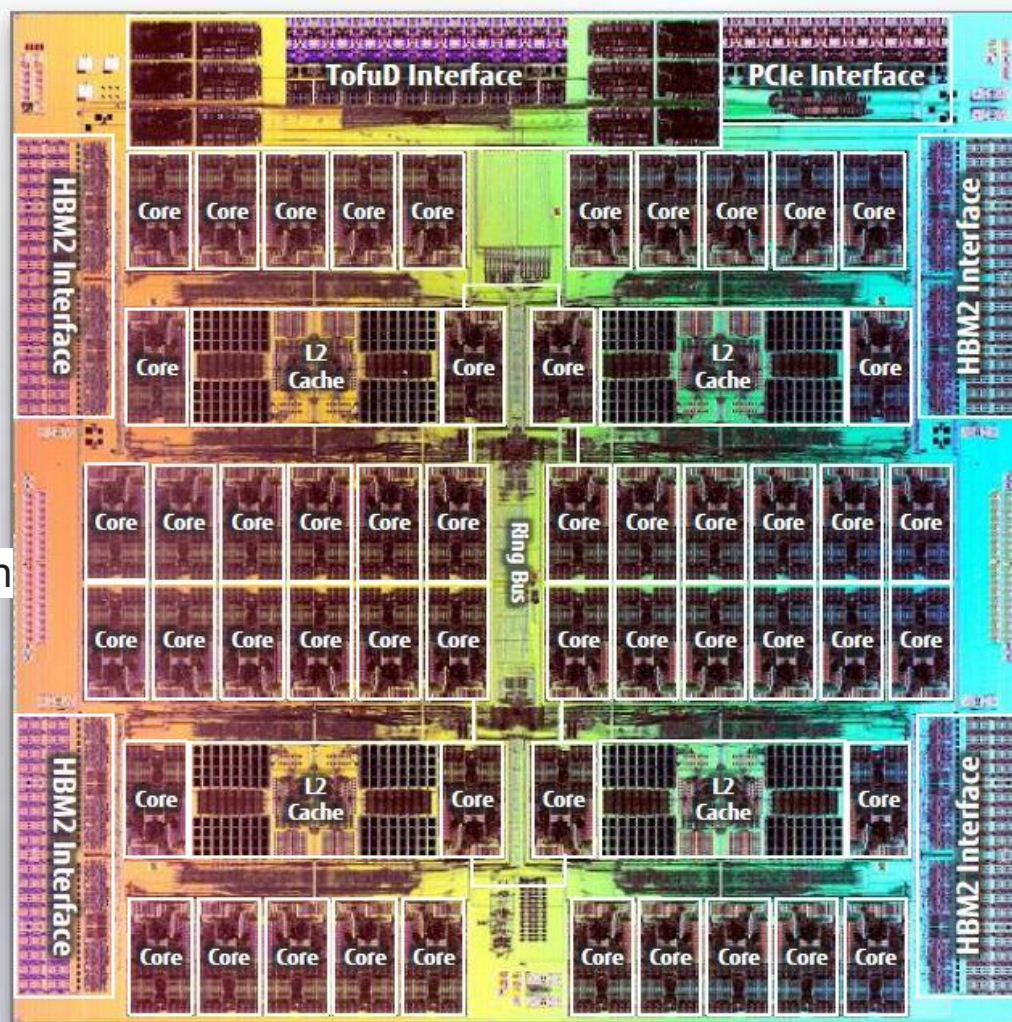
PRAM Model



A64FX
CPU
4 NUMA nodes
12 compute cores each

Memory : 32 GB

FLOPS : $13.5 \cdot 10^{12}$



4 performance cores

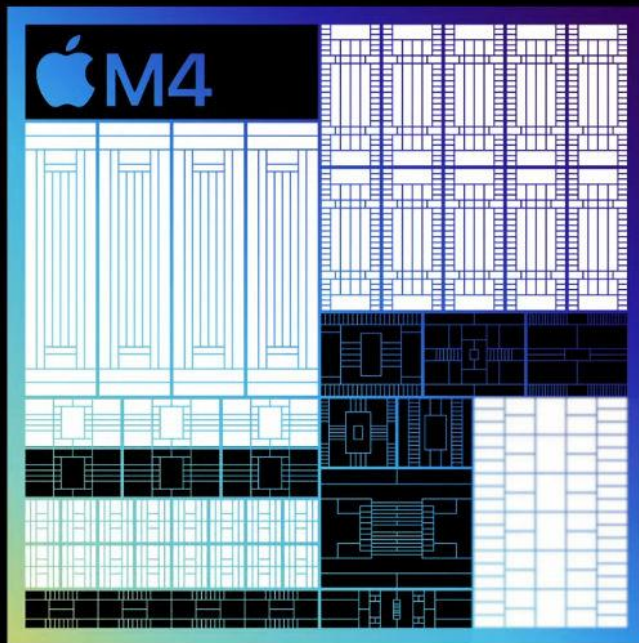
Improved branch prediction
Wider decode and execution engines
Next-generation ML accelerators

6 efficiency cores

Improved branch prediction
Deeper execution engine
Next-generation ML accelerators

Neural Engine

16-core design
Faster and more efficient



10-core GPU

Next-generation architecture
Dynamic Caching
Mesh shading
Ray tracing

Display engine

Tandem OLED support
Brightness and color compensation
10Hz-120Hz ProMotion support

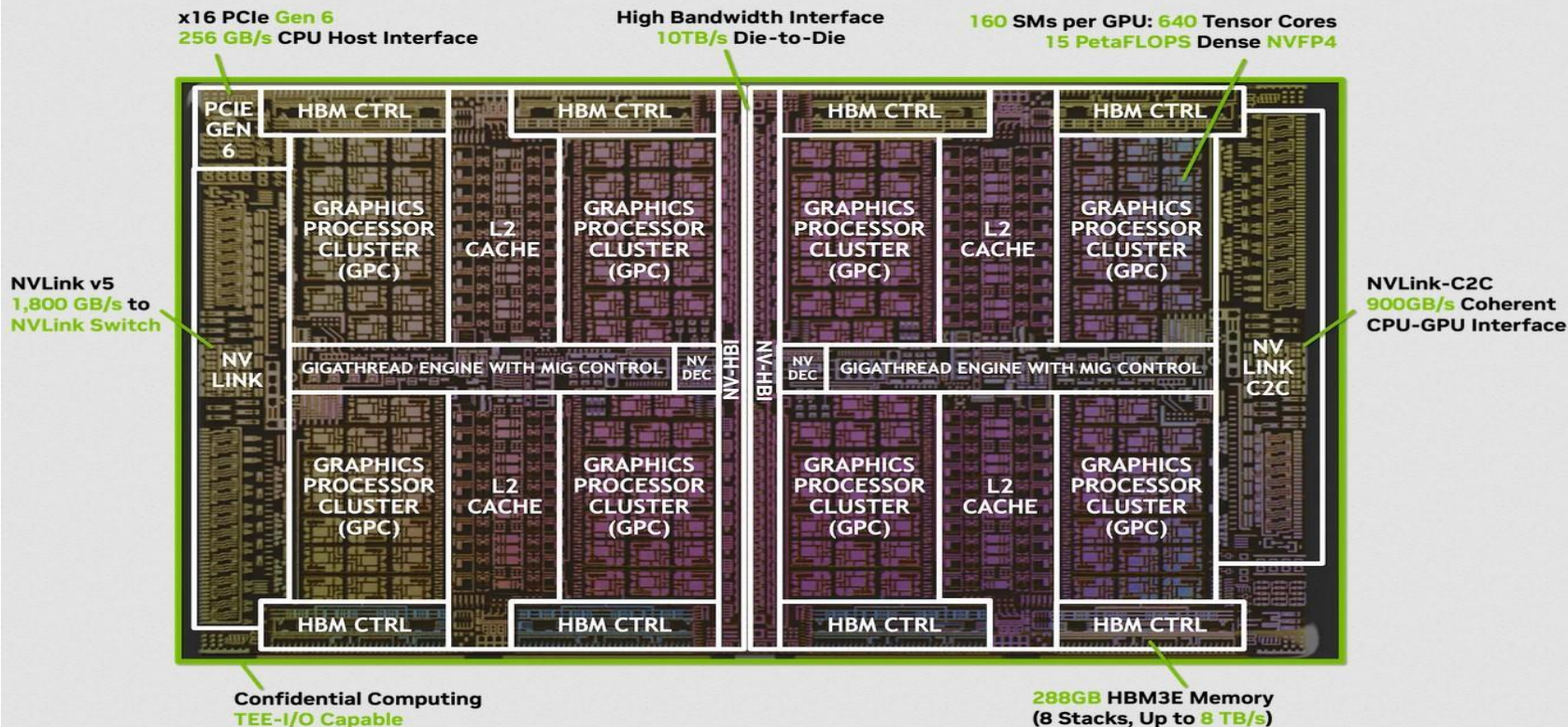
Apple M4

<https://www.apple.com/fr/newsroom/2024/05/apple-introduces-m4-chip/>

FP 16 : 8.52 TFLOPS
FP 32 : 4.26 TFLOPS

Execution units : 160
Max. Memory : 20GB

NVIDIA Blackwell Ultra GPU



Nvidia Blackwell

FP 16 : 2.25 PFLOPS
FP 32 : 1.12 PFLOPS

Execution units : 192
Max. Memory : 48GB

Serial computing should be re-invented

→ Realities of computing

- ◆ There are tons of wasted cycles
- ◆ CPU utilization typically <10% (of useful work)

→ Many other things limit performance

- ◆ Pipeline stalls
- ◆ Lock interference
- ◆ Waiting for I/O and network
- ◆ Data dependencies

→ Writing serial programs is broken

- ◆ Parallelism is everywhere
- ◆ Must exploit it to realize time efficiency, power savings

About parallelism

- Why do I **want** to write parallel programs?
 - ◆ to solve problems faster (strong scaling)
 - ◆ to solve bigger problems (weak scaling and memory)
- Why I **do not want** to write parallel programs?
 - ◆ tools are more difficult to use: expect 10x programming effort
 - ◆ for many problem performance does not improve

This course will help to decide when to develop a parallel approach and how to write it.

Context

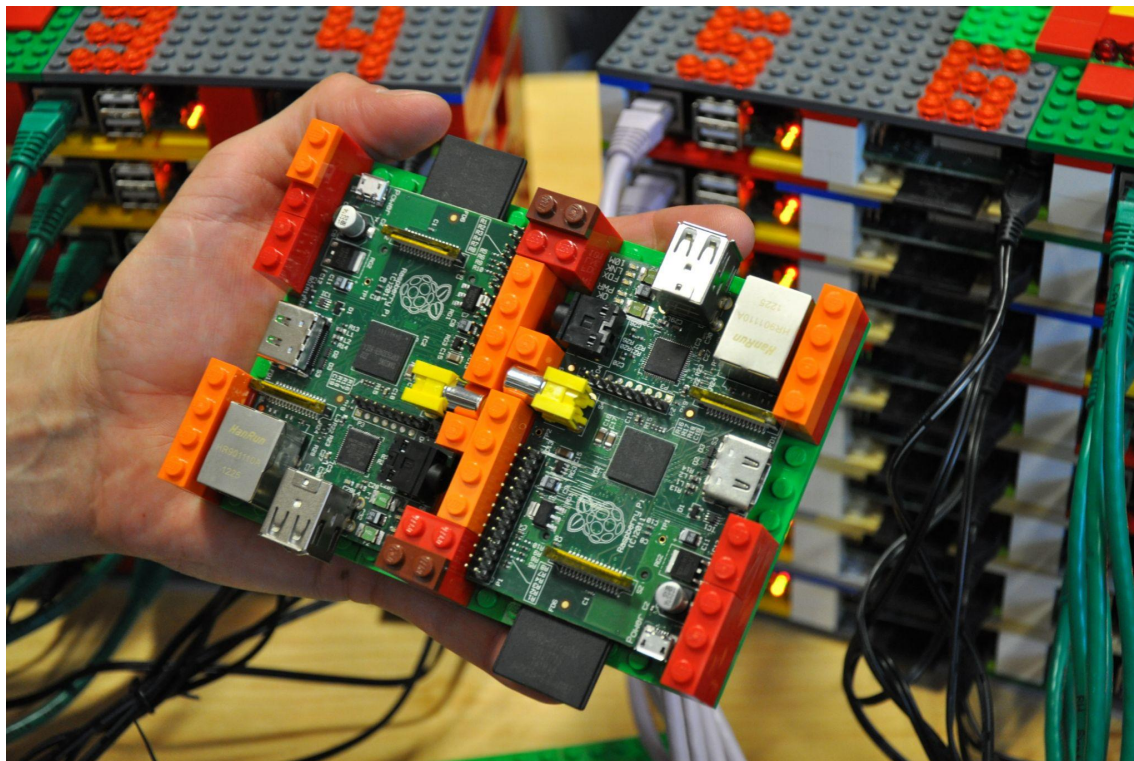
- Decrease time to solution.
- Solve larger problem.
- Combine resources of several processing units: gain access to more memory and more processing power.
- Harness the processing power of modern architectures.
- Use idle computer to perform embarrassing parallelism computation (SETI@home).
- Improve the precision of computations in a limited time (weather forecast).

How to achieve efficient parallelism?

- Processors: multicore, memory, network, accelerators, instructions.
- Compilers: dedicated library, automatic parallelism.
- Algorithms: tailored algorithms.
- Mathematics: adapted numerical methods, evolutionary methods.

Few words on hardware

A student HPC system

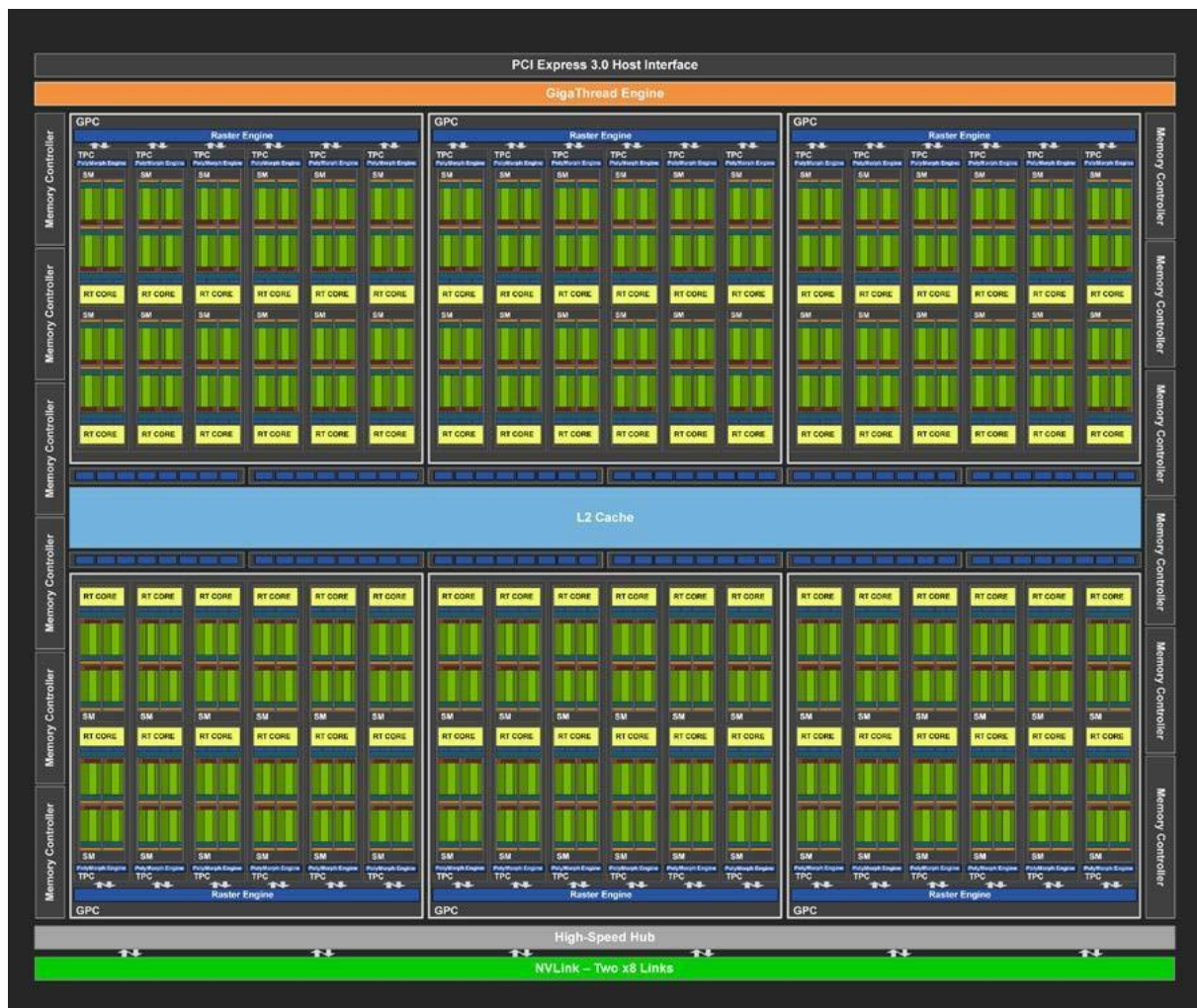


A “real” HPC system



What you will be working with!

- 33 RTX 6000 are available.
- Ressources are on virtual machine (up to 3 students on each card)
- Technical details :
 - ◆ CUDA Threads 4,608
 - ◆ NVIDIA Tensor threads 576
 - ◆ NVIDIA RT 72
 - ◆ GPU 24 GB GDDR6
 - ◆ Performances FP32 - 16,3 TFLOPS



Conclusions

Every computer is a parallel computer

- Parallel computers need independent work to run their many cores (or other resources) efficiently

Themes of this class

- Identify available parallelism in application
- Design parallel approaches
- Understand parallel hardware and how to optimize parallel performance