

CS110A

Introduction to Computer Systems

Chundong Wang

Course Info

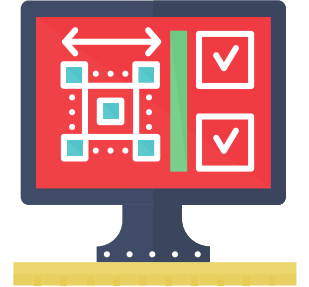
- Instructor and TA
 - WANG Chundong (email: wangc)
 - CUI Hanjia (email: cuihj2024)
- Website of CS110A
- Office Hours
 - 14:00 – 16:00, every Friday except public holidays



Outline

- History and backgrounds
 - Recent trends
- Course content
- Course rules/elements
- Great ideas in computer architecture

What is a computer?



- A **computer** is a machine that can be programmed to carry out **sequences of arithmetic or logical operations** (computation) **automatically**. Modern digital electronic computers can perform generic sets of operations known as **programs**. These programs enable computers to perform a wide range of tasks.
- A **computer system** is a nominally complete computer that includes the **hardware**, **operating system** (main software), and **peripheral equipment** needed and used for full operation. This term may also refer to a group of computers that are linked and function together, such as a computer network or computer cluster.

Early computers (calculators)

Chinese abacus



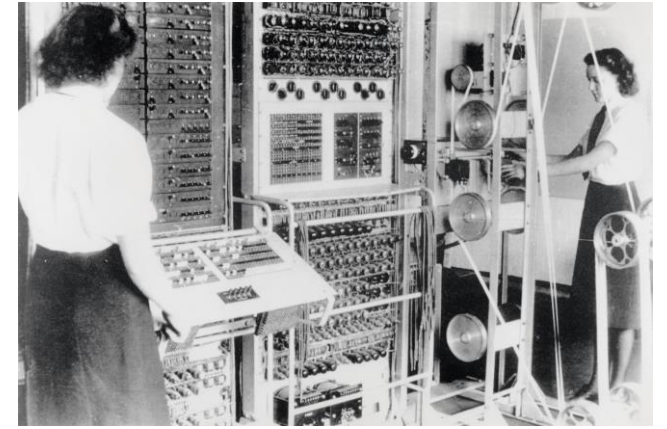
By Loadmaster (David R. Tribble). This image was made by Loadmaster (David R. Tribble). CC BY-SA 3.0.

Automatic mechanical calculator (1820)



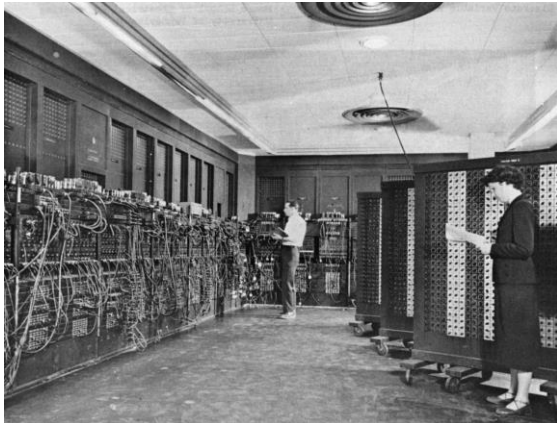
By Jitze Couperus from Los Altos Hills, California, USA. Uploaded by oxyman, CC BY 2.0.

Colossus (WWII, 1943-1945)



This file is from the collections of The National Archives (United Kingdom), catalogued under document record FO850/234.

ENIAC (1945-1946)



- Programmed by manually setting plugs and switches
- Vacuum tubes

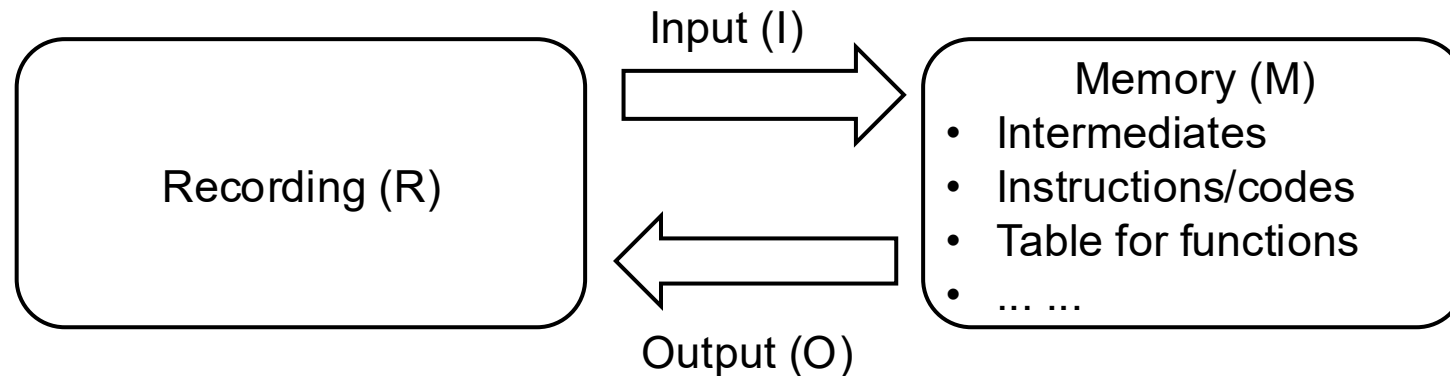
- 5000 additions or subtractions per second
- Modules to multiply, divide, and square root
- 30 tons, 200 kW, 18,000 vacuum tubes

Computer History

- Evolution of theory
 - Alan Turing proposed theoretical basis for the stored-program computer (1936)
 - *The First Draft of a Report on the EDVAC* by Von Neumann (1945)
 - Konrad Zuse suggested and implemented similar ideas (Harvard architecture)

<https://www.ibm.com/history/magnetic-tape>
<https://www.ibm.com/history/punched-card>

Back then punched cards and magnetic tapes

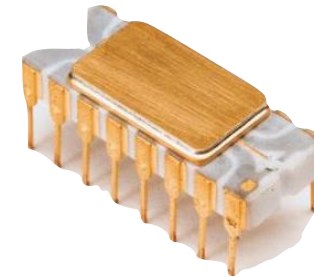


Computer History

- Evolution of theory
 - Alan Turing proposed theoretical basis for the stored-program computer (1936)
 - *The First Draft of a Report on the EDVAC* by Von Neumann (1945)
 - Konrad Zuse suggested and implemented similar ideas (Harvard architecture)
- Evolution of hardware & devices
 - Random-access digital storage (Williams tube, 1947)
 - Metal-oxide-silicon field-effect transistor (MOSFET, 1970-ish)
 - Integrated circuits (ICs) and then **VLSI**
 - Too many to list ...
- Microprocessor CPU (Intel 4004, 1971)



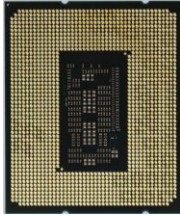
The first delivery of a fully operational 4004 was in March 1971 to Busicom for its 141-PF printing calculator engineering prototype (now displayed in the [Computer History Museum](#) in Mountain View, California).



From Intel

Modern Computers

Intel i7 12700



<https://maj191.com/product/intel-core-i7-12700-3-6ghz-cpu-25mb-cache-lga1700-tray/>
https://www.ocinside.de/review/intel_core_i7_12700k/3/

Processor

Cache

Control

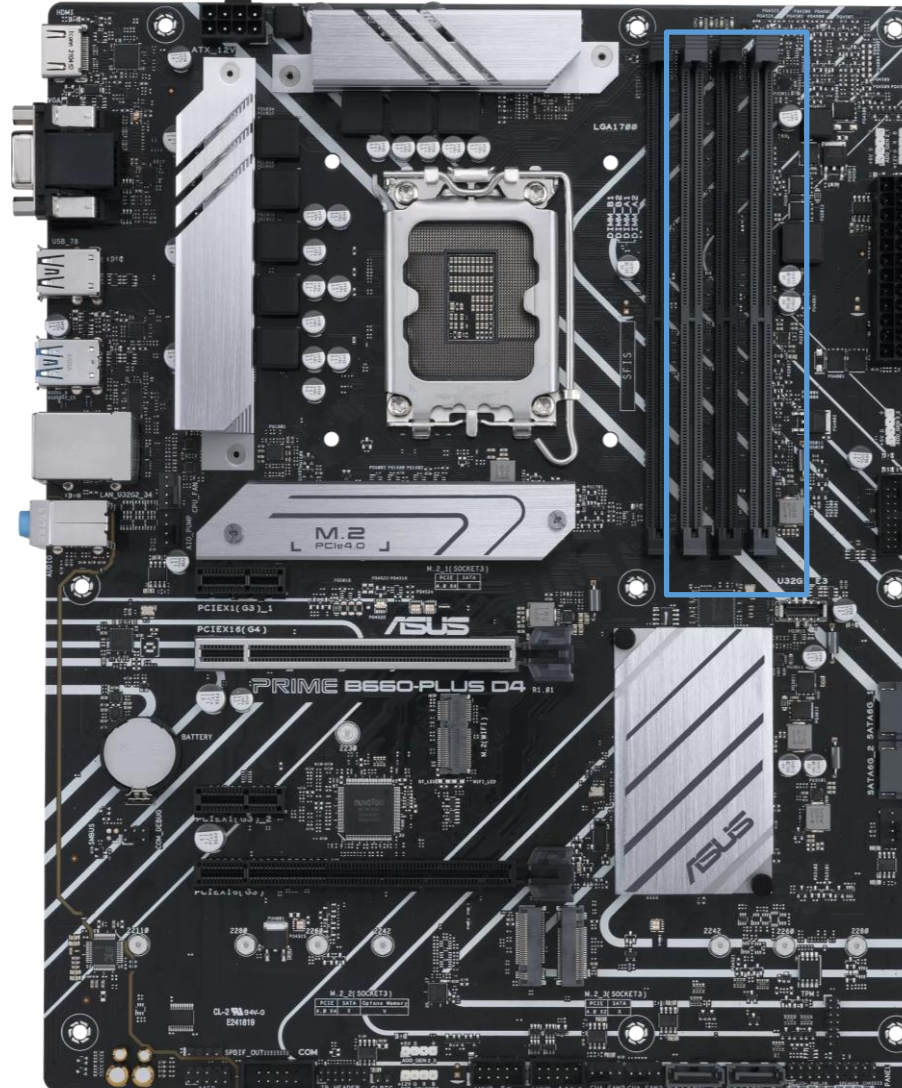
Datapath

PC

Registers

Arithmetic & Logic Unit (ALU)

PCB



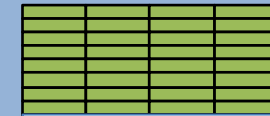
<https://www.asus.com/motherboards-components/motherboards/prime/prime-b660-plus-d4/>

DIMM
DDR4
5066 MHz

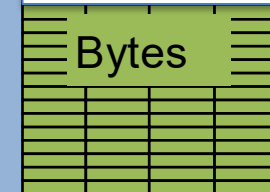
SDRAM



Memory



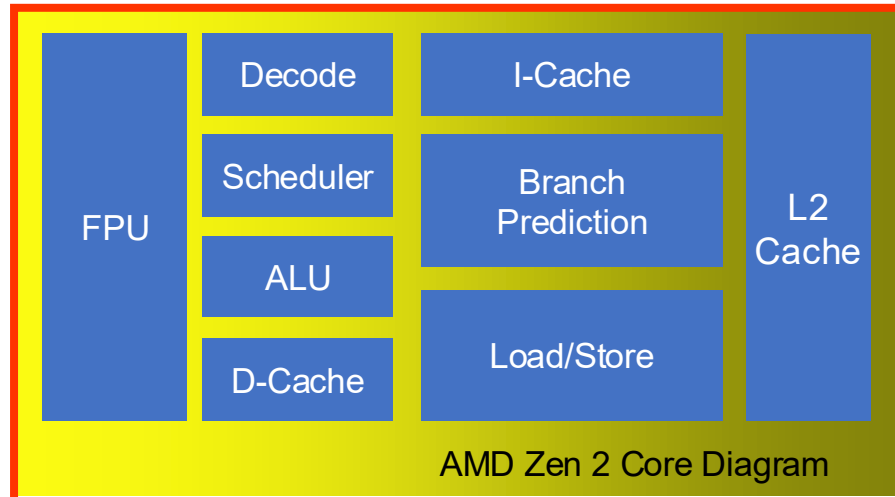
Program



Bytes

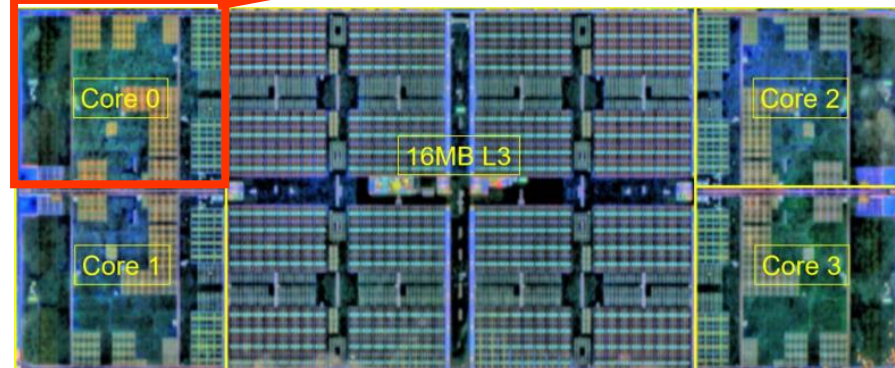
Data

Inside the CPU



Unit	Zen	Zen 2
Floating Point	128b	256b
L0 Branch Target Buffer	8 entries	16 entries
L1 Branch Target Buffer	256 entries	512 entries
L2 Branch Target Buffer	4K entries	7K entries
Op Cache	2K ops	4K ops
Integer Physical Register File	168 entries	180 entries
Integer Scheduler	84 entries	92 entries
AGEN	2	3
ROB	192 entries	224 entries
L2DTLB	1.5K	2K
L3 Cache Size	8MB	16MB

7.83 mm² per core



[1] T. Singh et al., "2.1 Zen 2: The AMD 7nm Energy-Efficient High-Performance x86-64 Microprocessor Core," IEEE International Solid- State Circuits Conference - (ISSCC), 2020, pp. 42-44.

Modern Computer Systems



From Huawei website

Inside an Intel Core i7 CPU

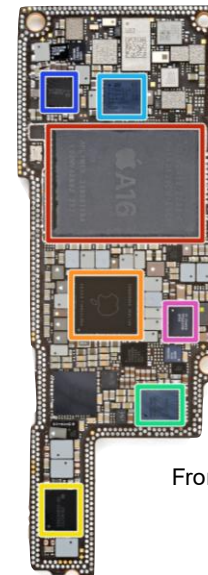
- 4 cores
- Up to 5.00 GHz
- Max 64 GB memory
- Include Intel Iris Xe Graphics (GPU)
- Windows/Linux (for desktop)



Inside an Apple A16 processor (SoC)

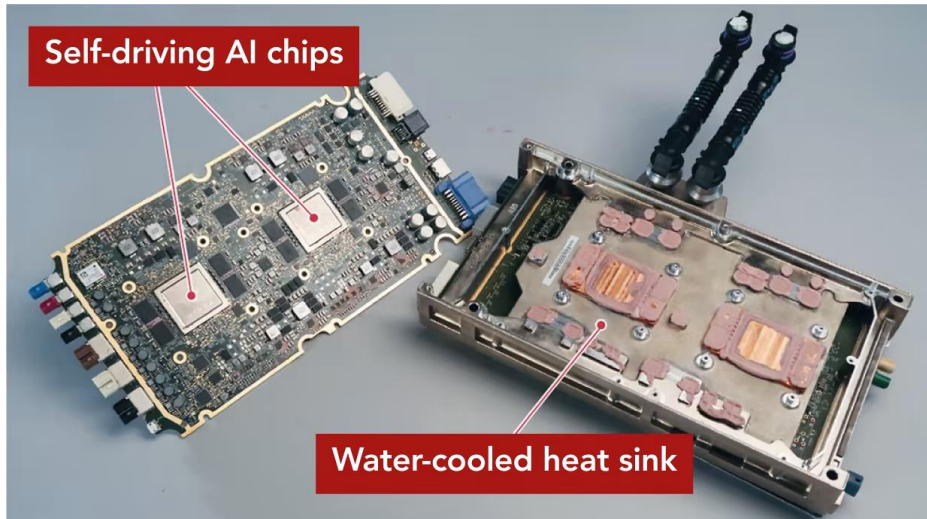
- 6 cores
- Up to 3.46 GHz
- 6 GB memory
- Include Apple-designed GPU, image processor, 16-core neural engine, etc.
- 17 TOPS on neural engine
- iOS

From Wikipedia



From ifixit.

Modern Computer (Embedded) Systems



FSD computer, inside a FSD chip (SoC)

- 12 CPU cores, 2.2 GHz
- Up to 8 GB memory
- Include a Mali GPU @ 1 GHz, 2 neural processing units @ 2 GHz, etc.
- GPU 600 GFLOPS, NPU 36.86 TOPS
- Specialized software

From Wikichip.



Main screen (infotainment) computer

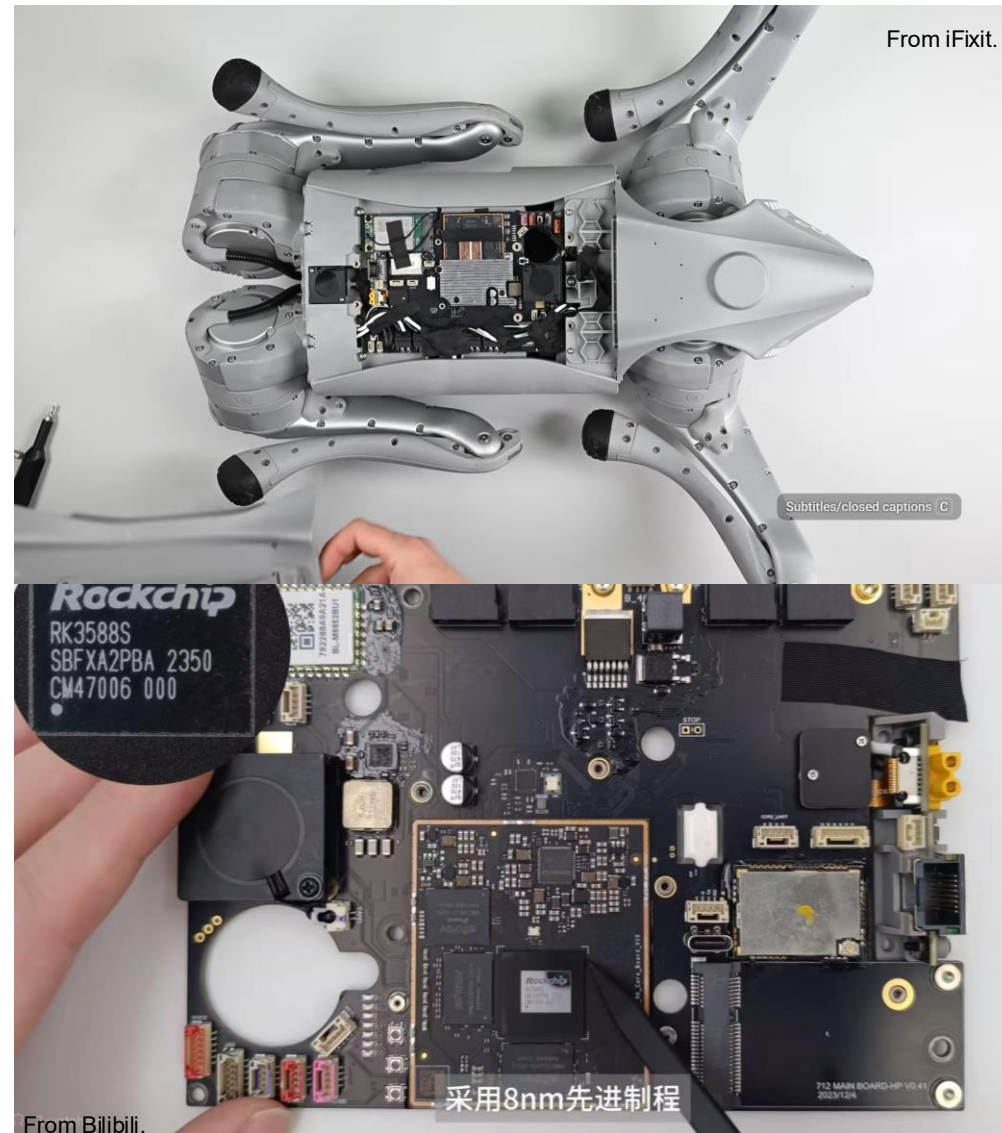
- Ubuntu-based OS
- AMD Ryzen & GPU/Intel Atom

From [Nikkei-xTech](https://www.nikkei-xTech.com), [Engineerix](https://www.engineerix.com) & [Dongchedi](https://www.dongchedi.com).
<https://www.youtube.com/watch?v=ODdIRr5RzI8>
<https://www.dongchedi.com/article/7122294255204712972>

Modern Computer (Embedded) Systems



From Unitree.



From Bilibili.

Modern Computer Systems



SIST computing center

- Hundreds/thousands of servers
- Each has one or multiple CPUs
- Mostly runs Linux OS

Hardware Revolution (Cloud Services)

- Heterogeneous

Providers	CPU	GPU	FPGA*	ASIC (DSA)
Alibaba	X86/ARM/RISC-V	Nvidia/AMD	Altera/AMD	AliNPU
AWS (Amazon)	X86/Graviton (ARM)	Nvidia/AMD	AMD	Trainium
Azure (MS)	X86	Nvidia	Altera	N/A
Baidu	X86	Nvidia	AMD	Kunlun
Google	X86	Nvidia	N/A	TPU
Huawei	X86/Kunpeng (ARM)	Nvidia & Ascend	AMD	Ascend
Tencent	X86	Nvidia & Xinghai	AMD	Enflame (燧原)

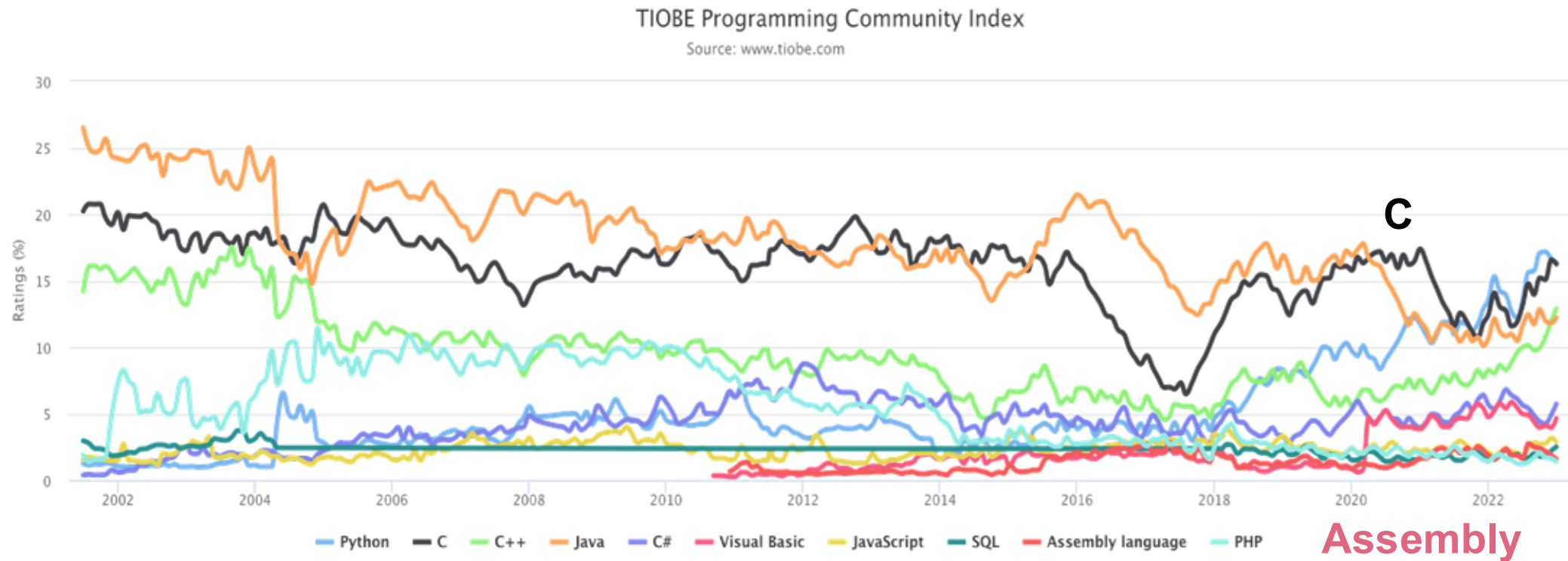
*AMD: formerly Xilinx

Computer Hardware Evolution

- Stronger (process complex information, e.g., AI)
 - High performance, high energy efficiency, etc.
- Easier to use (touch screen, voice control, etc.)
- Diversity (edge/cloud, various architectures)
- Heterogeneous (CPU/GPU/ASIC/FPGA)
- All developed based on computer architecture theory

Programming (Software)

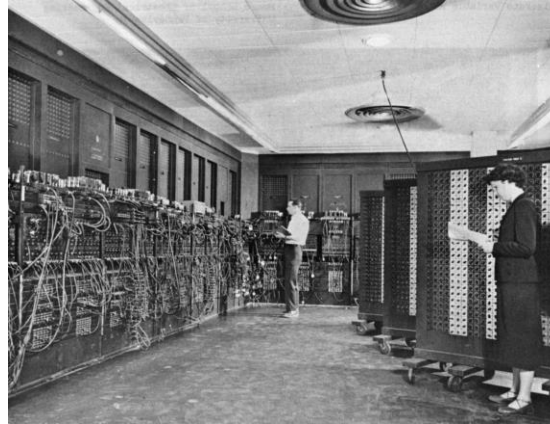
- Write program (as CS students/programmers)



Programming Evolution



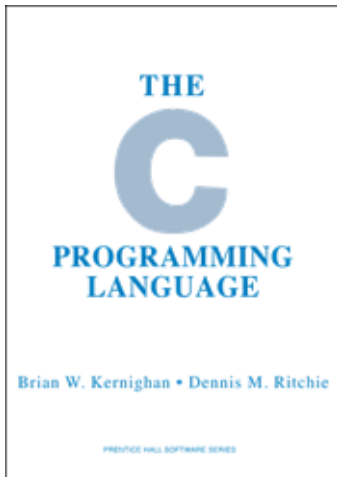
Ada Lovelace,
concept of
program



ENIAC, switches and relays



Margaret Hamilton with
the (**assembly**) code she
wrote.



General purpose, procedural
programming language

C



Domain or “hardware” specific

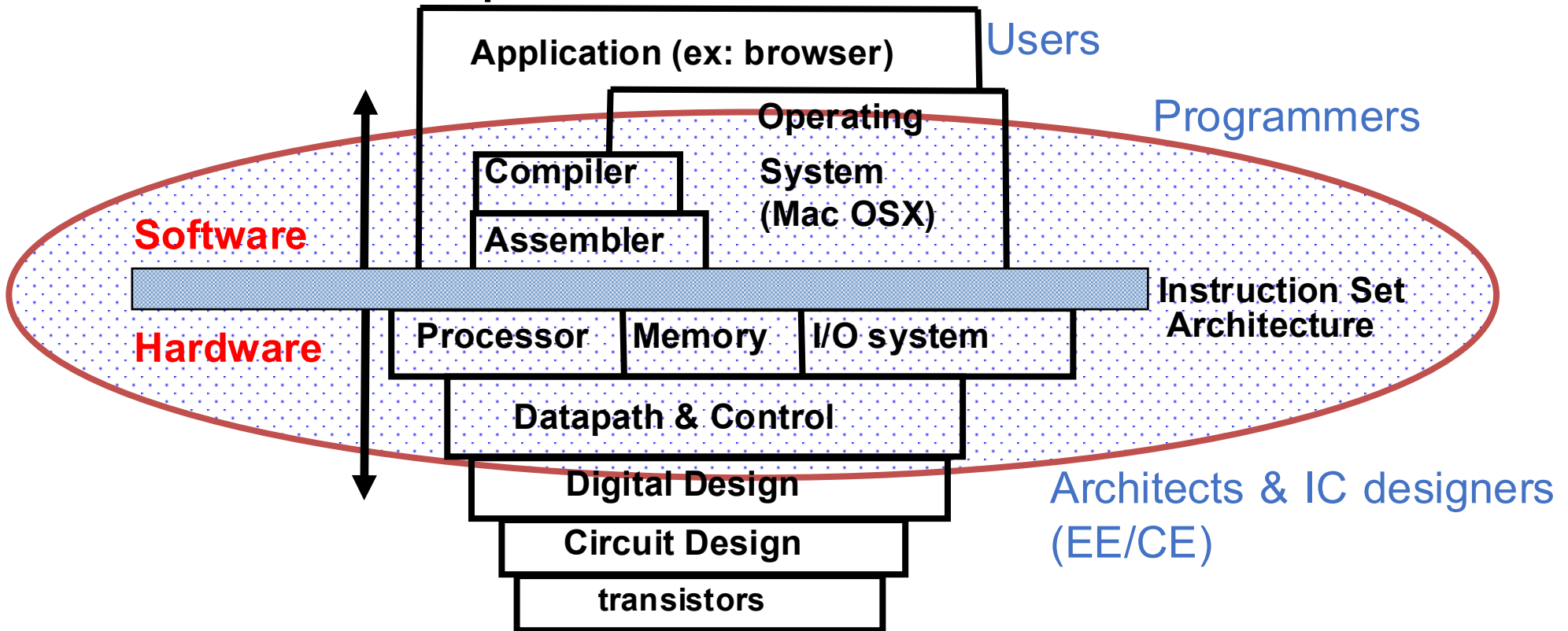


A new era

“It is our job to create computing technology such that nobody has to program” by Jensen Huang.

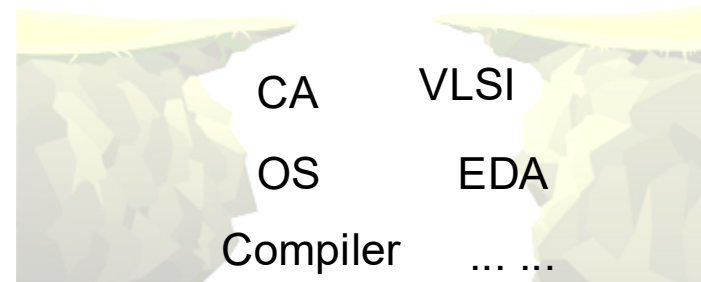
CS110A is TO EXPLAIN HOW COMPUTER WORKS (FUNDAMENTALS)

- Abstraction of computers

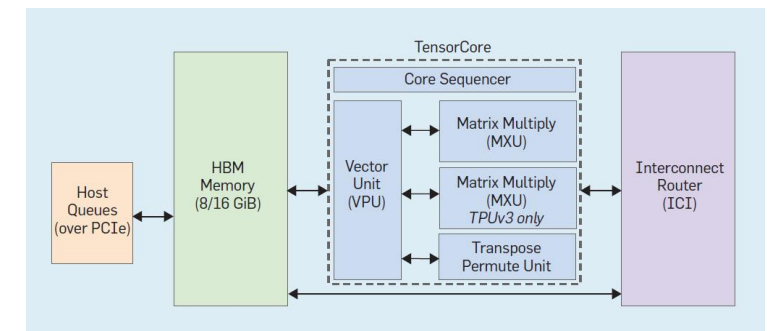
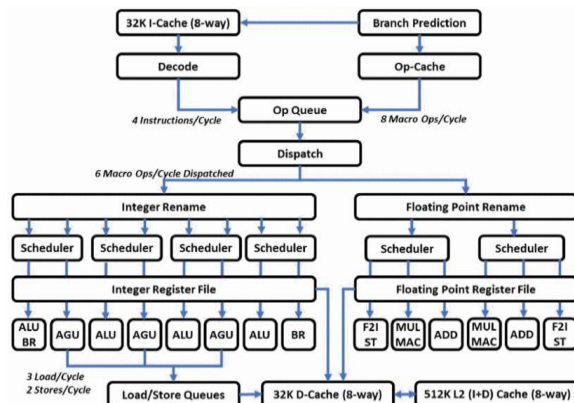


CS110A is TO EXPLAIN HOW COMPUTER WORKS (FUNDAMENTALS)

- It is about the hardware-software interface
 - How is the high-level language translated to machine code
 - How to build the hardware (digital circuit) to understand the machine code and execute corresponding instructions (CPU/MCU design)



<https://www.partsocarengine.com/parts-of-a-car/>



CS110A is TO EXPLAIN HOW COMPUTER WORKS (FUNDAMENTALS)

- It is about the **hardware-software** interface
 - How is the high-level language translated to machine code
 - How to build the hardware (digital circuit) to understand the machine code and execute corresponding instructions (CPU/MCU design)
- We will learn
 - Information representation
 - C review and memory management
 - CALL (compiler, assembler, linker & loader, basics)
 - ISA & assembly (RISC-V as example)
 - Digital circuit design

CS110A is TO TRAIN PROFESSIONALS

- It is about the **hardware-software** interface
 - How is the high-level language translated to machine code
 - How to build the hardware (digital circuit) to understand the machine code and execute corresponding instructions (CPU/MCU design)
 - Evaluate **tradeoffs** of different designs and ideas
- How to improve the performance of computers (parallelism/memory hierarchy)
- What does the programmer need to know to achieve the highest possible performance



CS110A is **NOT** really about Programming

- Programming language as Tools
 - **C** is close to the underlying hardware, unlike languages like Python, Java!
 - Allows us to talk about key hardware features in higher level terms
 - Allows programmer to explicitly harness underlying hardware parallelism for **higher performance** and **power efficiency**
 - We will also learn **assembly**
 - Allows us to talk about the hardware in lower level terms
 - Uses **RISC-V** assembly/ISA as an example



Parallel Thread eXecution

PTX ISA

CS110A is **to connect system with AI**

- The development of systems boosts the development of AI
- AI Infrastructures
 - The physical and software systems used to develop, train, deploy, and operate AI models and applications.
- Classic system ideas are useful for AI
- AI in turn helps the development of systems

Course contents

- How hardware & program orchestrate in a computer?
 - Software side:
 - Information representation
 - C review and memory management
 - Assembly (RISC-V ISA)
 - Compiler, assembler, linker & loader (in general)
 - CPU hardware design & How CPU hardware works (RISC-V ISA)
 - Logic & synchronous digital systems
 - Functional Units, FSM & Datapath
 - Pipeline, superscalar (Parallelism 1, Instruction-level parallelism)
 - Multi-issue
- Memory systems (Cache, Main memory)
- Parallelism (SIMD, TLP, Sync & OpenMP, data/thread-level parallelism)
- System topics (OS, interrupts, I/O, virtual memory, FPGA, distributed computing, DMA, external memory, fault-tolerance, security, etc.)
- AI related topic (in the system's perspective)

Why These Topics?

Because you are supposed to be a computer scientist.

- Regardless of your future direction, learning the principles of digital design & computer architecture will be useful to
 - design better hardware (computer engineering, IC design)
 - design better software (as programmer)
 - design better systems (as architects)
 - make better tradeoffs in design (required for all CS-related tasks)
 - understand why computers behave the way they do
 - solve problems better
 - think “in parallel”
 - to resolve AI-related problems as a systems architect
 - and so on

A comparison between CS110 and CS110A

Similarity

- Both on architecture and systems, sharing many contents.
- Both being fundamental and essential for CS students.
- Both in the system's perspective.
- Both with assignments and exams.
- Both very useful, hopefully.

Difference

- CS110 has an affiliated course named CS110P, CS110A no.
- CS110 has in-person labs, CS110A no.
- CS110 is not so “AI”, CS110A yes.
- Many topics covered in CS110A, not in CS110.
 - Vice versa.
- Fewer workloads in CS110A, *hopefully*.

CS110A grading

- Homework (20%)
- Course projects (24%)
- Exams
 - Mid-term exam (24%)
 - Final exam (30%)
- Participation and reward (2%+)

Policy on Assignments and Independent Work

- ALL LABS/HOMEWORKS/ASSIGNMENTS, AND PROJECTS WILL BE DONE INDEPENDENTLY
- With the exception of laboratories and assignments that explicitly permit you to work in groups, all homework are to be YOUR work and your work ALONE.
- PARTNER TEAMS MAY NOT WORK WITH OTHER PARTNER TEAMS
- You can discuss your assignments with other students, and credit will be assigned to students who help others by answering questions on Piazza (participation), but we expect that what you hand in is yours.
- Level of detail allowed to discuss with other students: Concepts (Material taught in the class/in the text book)! [Pseudocode or code is NOT allowed!](#)
- Use the Office Hours of the TA and the Profs. if you need help with your homework/project!
- [Rather submit an incomplete homework with maybe 0 points than risking an F!](#)
- It is NOT acceptable to copy solutions/code from other sources (other students/web/**AI-generated**).
- You can never look at homework/ project code not by you/ your team!
- You cannot give your code to anybody else → secure your computer when not around it
- [It is NOT acceptable to copy \(or start your\) solutions from the Web/other students/AI-generated.](#)
- [It is NOT acceptable to use PUBLIC github archives \(giving your answers away\)](#)
- [It is NOT acceptable to give anyone other than your project partner access to your gitlab/github/...!](#)
- [Protect your code and password! Do not allow other students peeping at your screen.](#)
- We have tools and methods, developed over many years, for detecting this. You WILL be caught, and the penalties WILL be severe.
- At the [minimum F](#) in the course, and [a letter to your university record](#) documenting the incidence of cheating.
- **Both Giver and Receiver are equally culpable and suffer equal penalties**

Any academic misbehavior will immediately fail the course.

All records will be reported to SIST, Office of Undergraduate Programs, and colleges.

If you ignore this warning, you will bear the consequences.

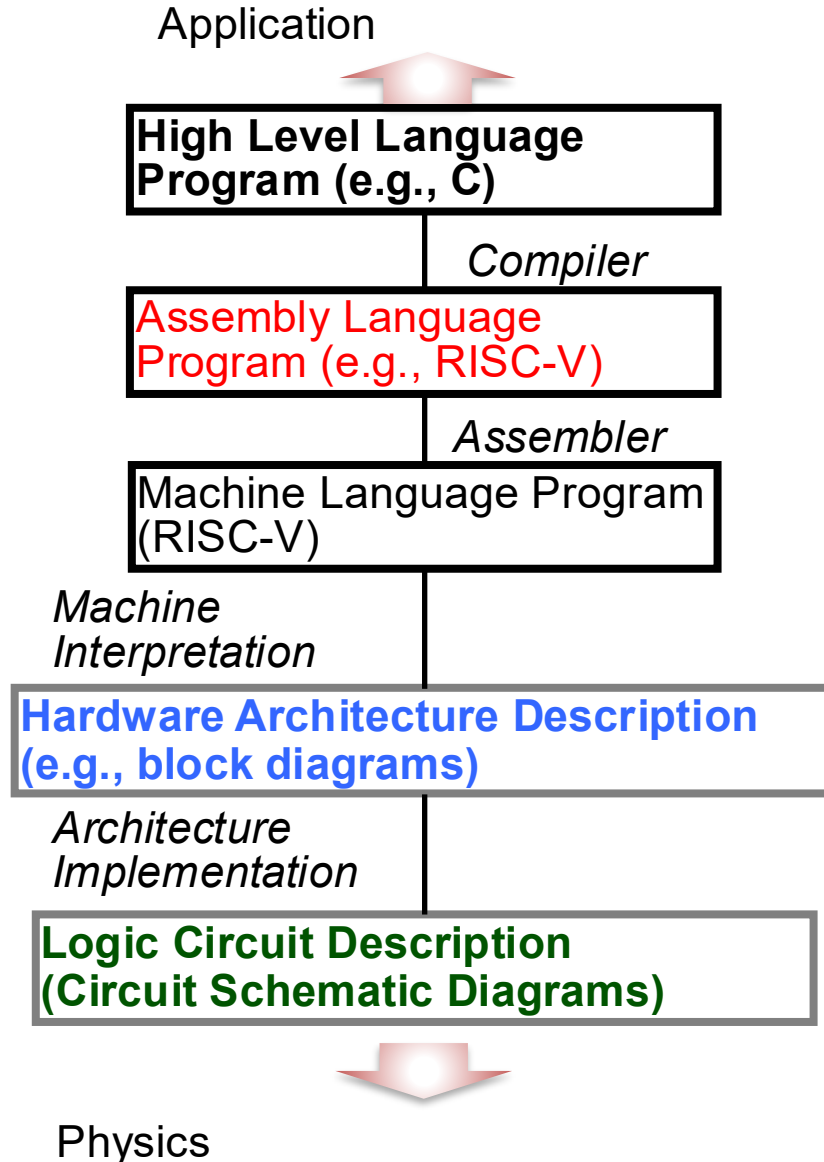
Great ideas in Computer Architecture

Great ideas

- Abstraction (layers of representation/interpretation)
- Moore's law (designing through trends)
- Make the common case fast
- Principle of locality (memory hierarchy)
- Parallelism (pipeline as a special case)
- Performance measurement & improvement
- Dependability via redundancy

Great Idea 1: Abstraction

(Levels of representation/Interpretation)

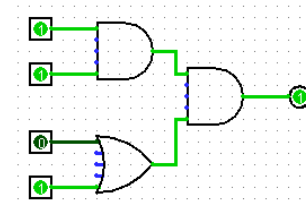
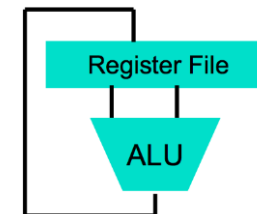


```
temp = v[k];  
v[k] = v[k+1];  
v[k+1] = temp;
```

```
lw  t0, 0(s2)  
lw  t1, 4(s2)  
sw  t1, 0(s2)  
sw  t0, 4(s2)
```

Anything can be represented
as a *number*,
i.e., data or instructions

```
0000 1001 1100 0110 1010 1111 0101 1000  
1010 1111 0101 1000 0000 1001 1100 0110  
1100 0110 1010 1111 0101 1000 0000 1001  
0101 1000 0000 1001 1100 0110 1010 1111
```



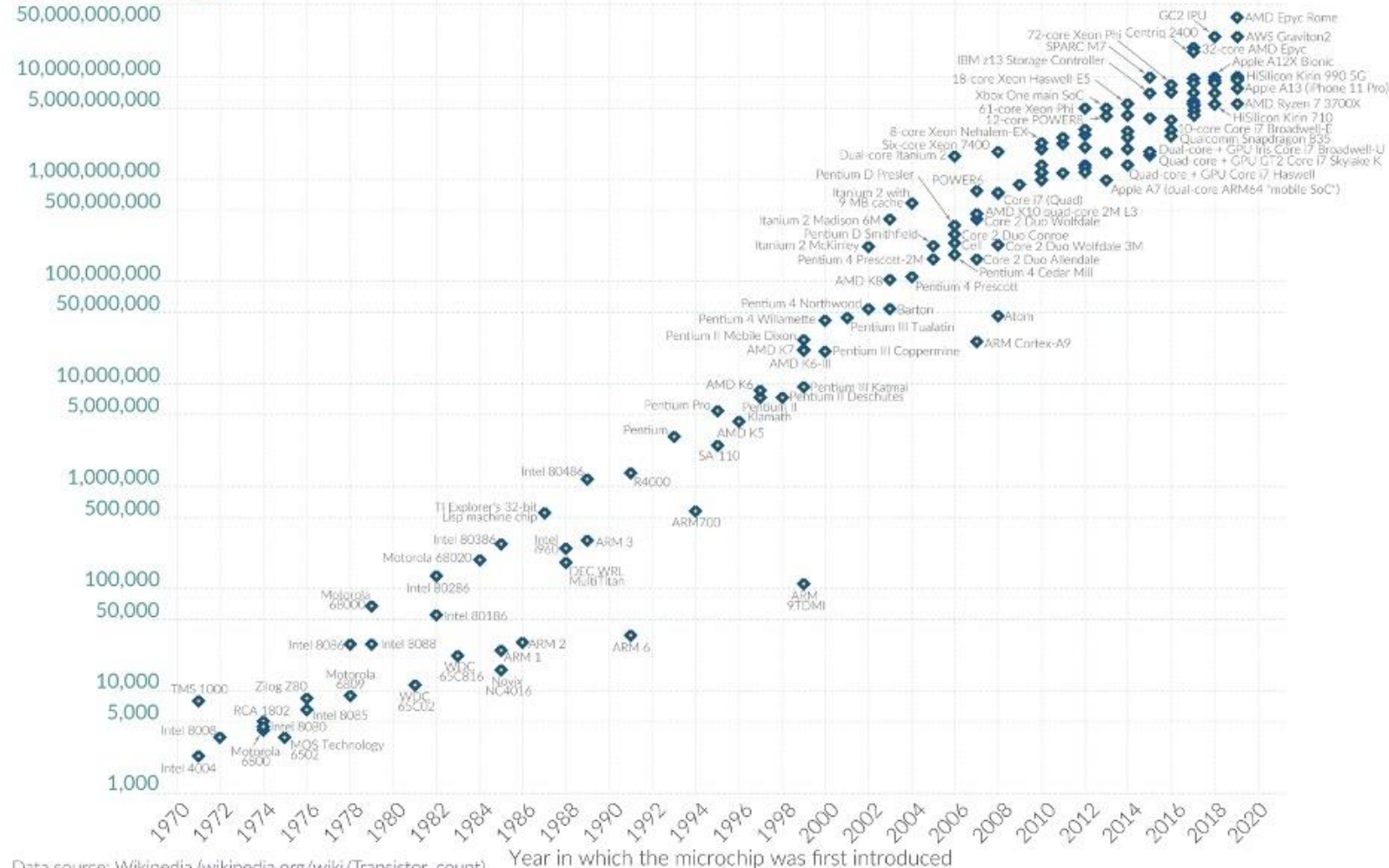
2. Moore's Law?

Moore's Law: The number of transistors on microchips doubles every two years

Moore's law describes the empirical regularity that the number of transistors on integrated circuits doubles approximately every two years. This advancement is important for other aspects of technological progress in computing – such as processing speed or the price of computers.

Our World
in Data

Transistor count

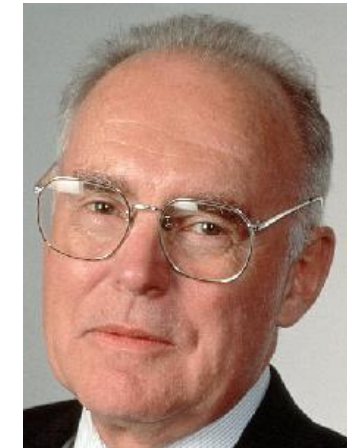


Data source: Wikipedia (wikipedia.org/wiki/Transistor_count)

OurWorldinData.org – Research and data to make progress against the world's largest problems.

Licensed under CC-BY by the authors Hannah Ritchie and Max Roser.

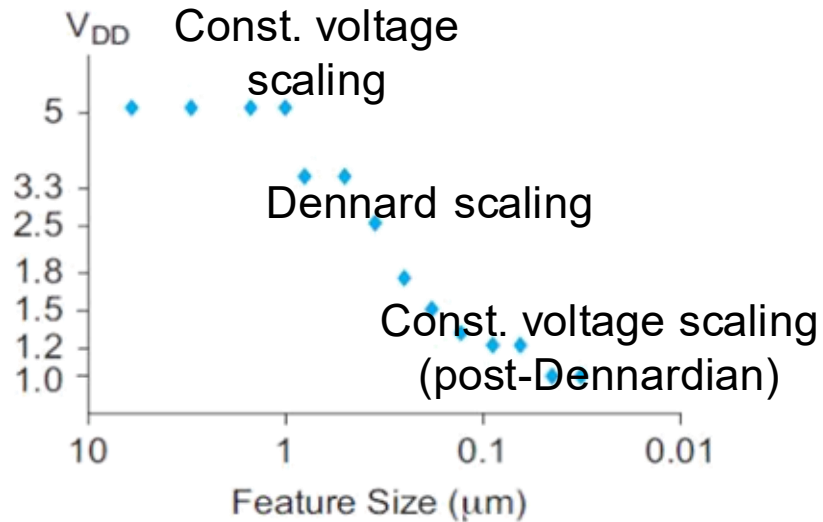
**Predict based on observation:
2X Transistors / chip
every 2 years (18 months)**



**Gordon Earle Moore
(1929 - 2023)
Intel Cofounder**

Interesting Times

- Dennard scaling (TDP)

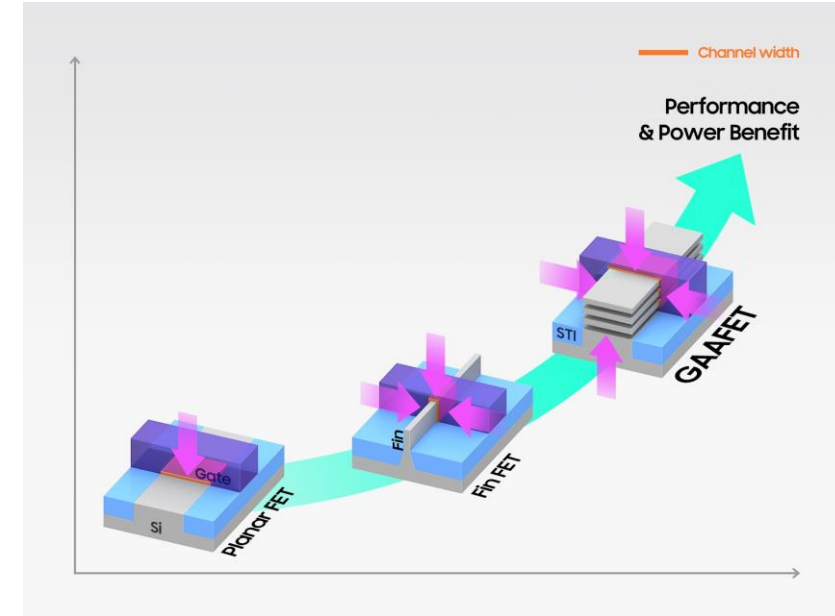
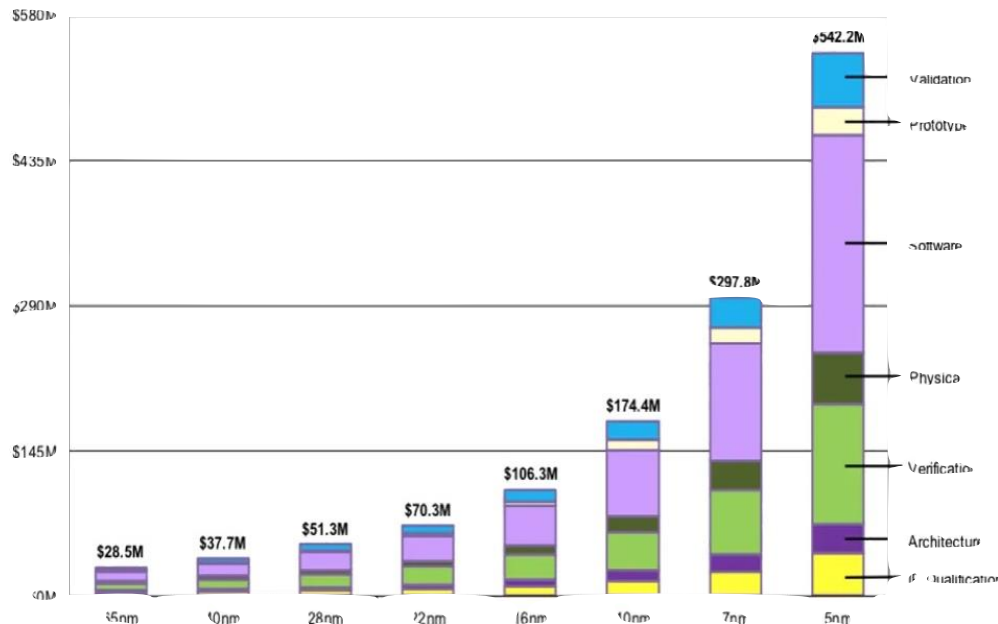


- Moore's Law relied on the cost of transistors scaling down as technology scaled to smaller and smaller feature sizes.
- BUT newest, smallest fabrication processes <5nm, might have greater cost/transistor !!!!
- So, why shrink???



Advanced Technology Nodes

- Now: 4 nm process (A16 Bionic, TSMC)
- Future: up to 3 nm (rumors on A17)
- Other notable players: Samsung, Intel, GlobalFoundries (AMD), IBM
- China: SMIC: 14-nm production



From Samsung Semiconductor

Design cost at different technology nodes.

Data source: Handel Jones, IBS, 2018.

3. Make Common Case Faster: Amdahl's Law

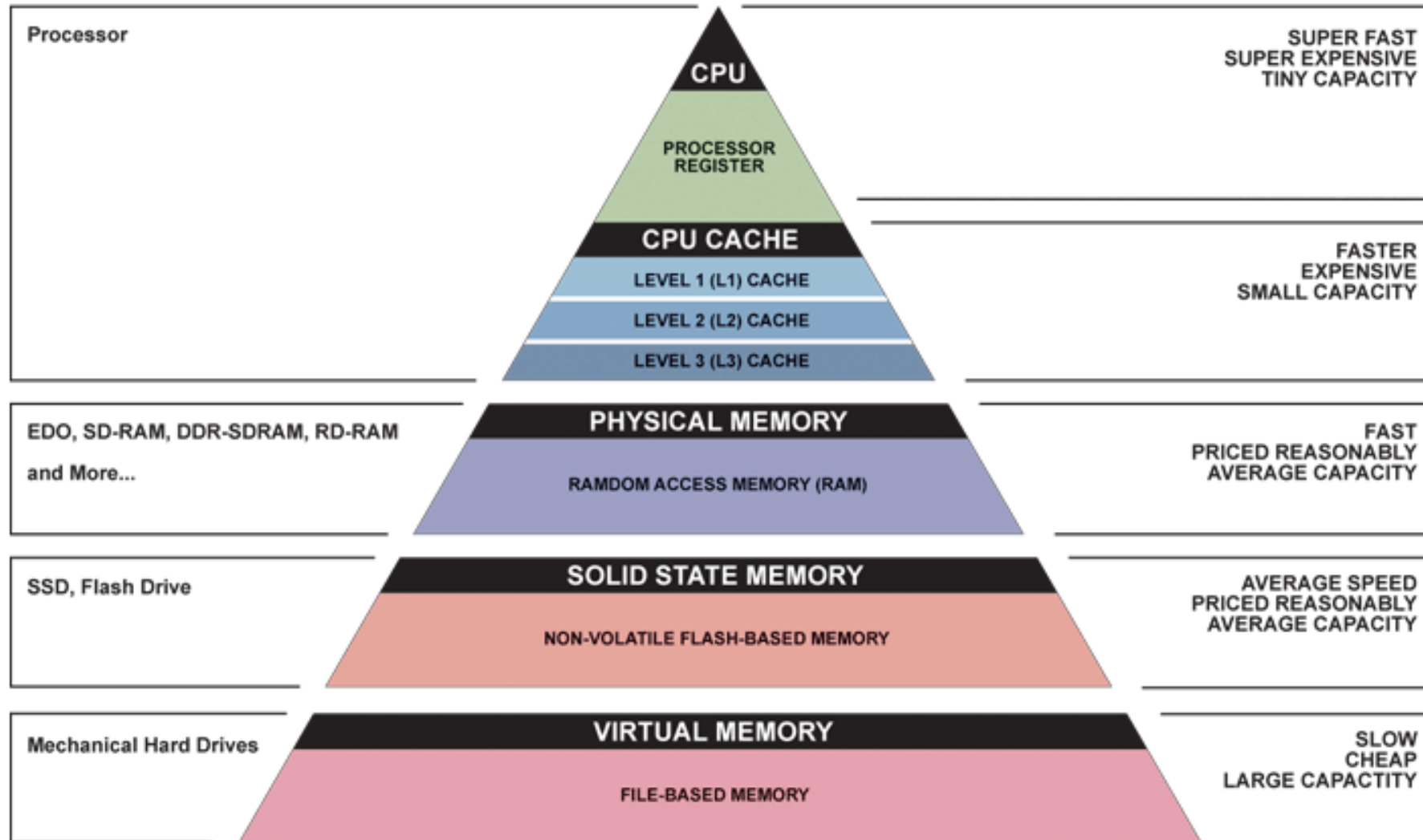
A rule stating that the performance enhancement possible with a given improvement is limited by the amount that the improved feature is used. It is a quantitative version of the **law of diminishing returns**.



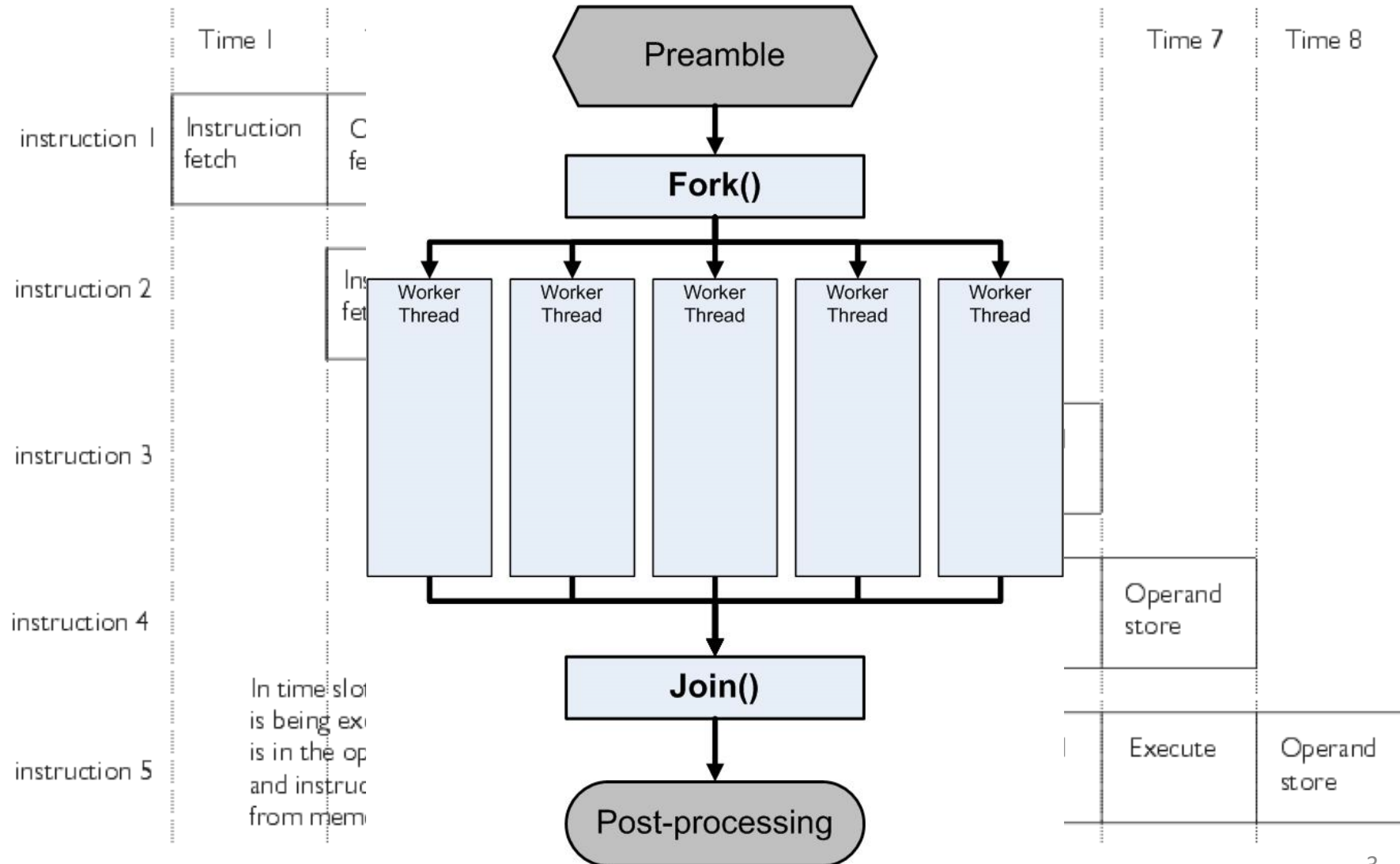
Gene Amdahl
Computer Pioneer

$$\begin{aligned} & \text{Execution time after improvement} \\ = & \frac{\text{Execution time affected by improvement}}{\text{Amount of improvement}} + \text{Execution time not affected} \\ & \text{(+ Overhead for implementing the improvement)} \end{aligned}$$

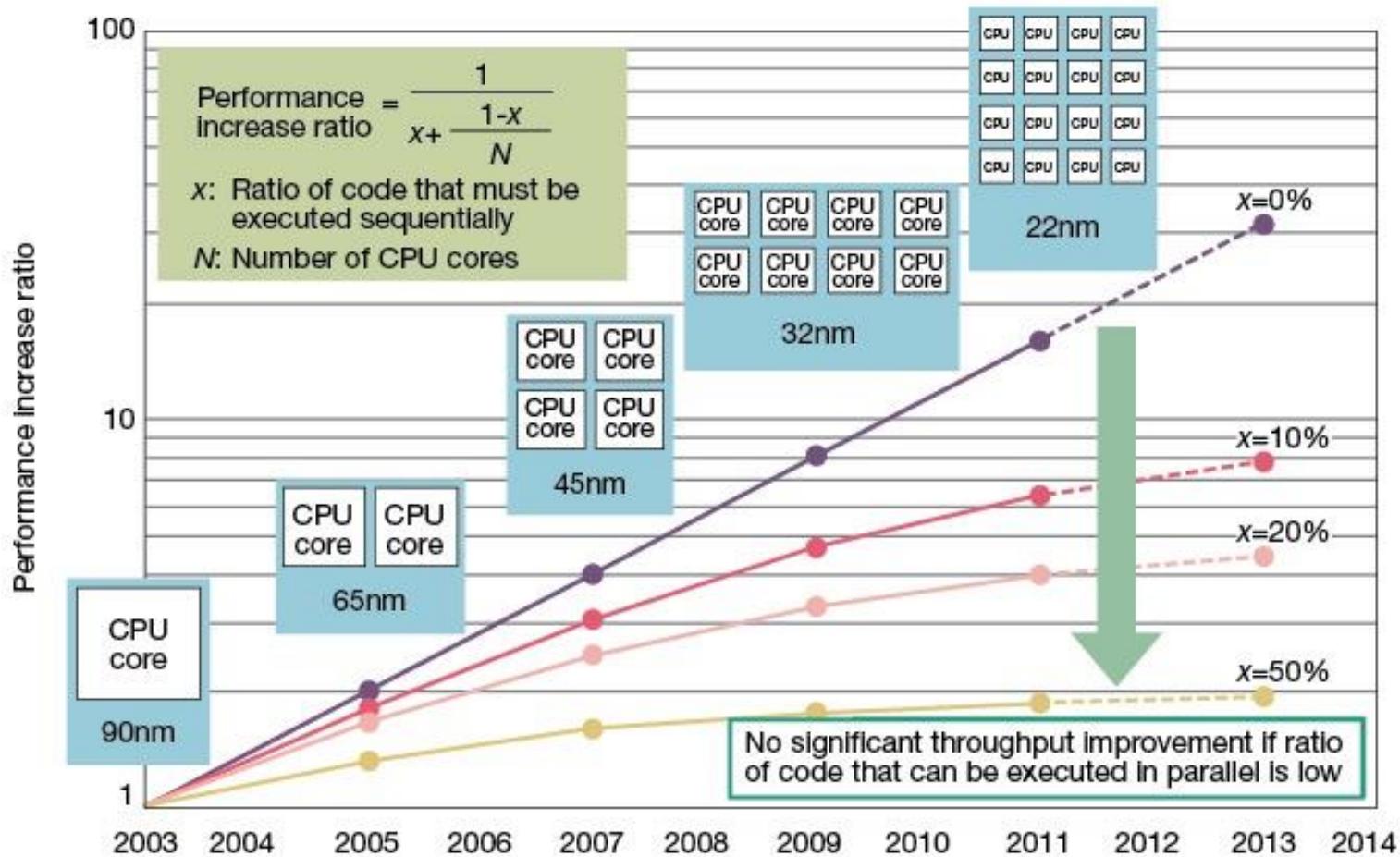
Great Idea 4: Principle of Locality/ Memory Hierarchy



Great Idea 5: Parallelism/Pipeline



Caveat: Amdahl's Law



Gene Amdahl
Computer Pioneer

Fig 3 Amdahl's Law an Obstacle to Improved Performance Performance will not rise in the same proportion as the increase in CPU cores. Performance gains are limited by the ratio of software processing that must be executed sequentially. Amdahl's Law is a major obstacle in boosting multicore microprocessor performance. Diagram assumes no overhead in parallel processing. Years shown for design rules based on Intel planned and actual technology. Core count assumed to double for each rule generation.

Great Idea 6: Performance Measurement & Improvement

- Latency
 - How long to set the problem up. It is all about time to finish.
- Power/energy consumption
 - How much energy consumed to perform certain tasks.
- Benchmark (SPEC 2023,
<https://www.top500.org/>, LINPACK)

Coping with Failures

- 4 disks/server, 50,000 servers
 - Assume 4% annual failure rate
- On average, how often does a disk fail?
 - 1 / month
 - 1 / week
 - 1 / day
 - 1 / hour

$50,000 \times 4 = 200,000$ disks

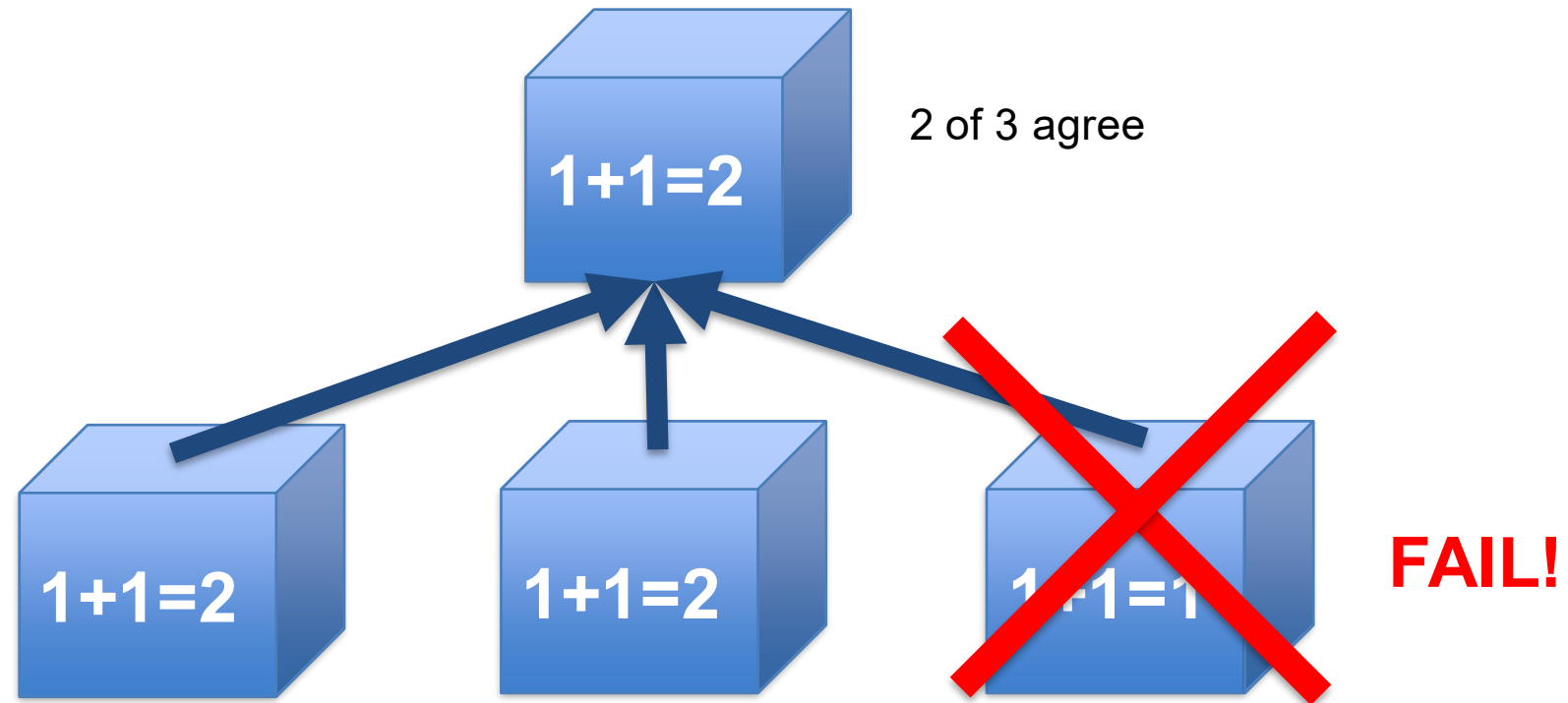
$200,000 \times 4\% = 8000$ disks fail

$365 \text{ days} \times 24 \text{ hours} = 8760$ hours

Great Idea 7:

Dependability via Redundancy

- Redundancy so that a failing piece doesn't make the whole system fails



Increasing transistor density reduces the cost of redundancy

Great Idea 7:

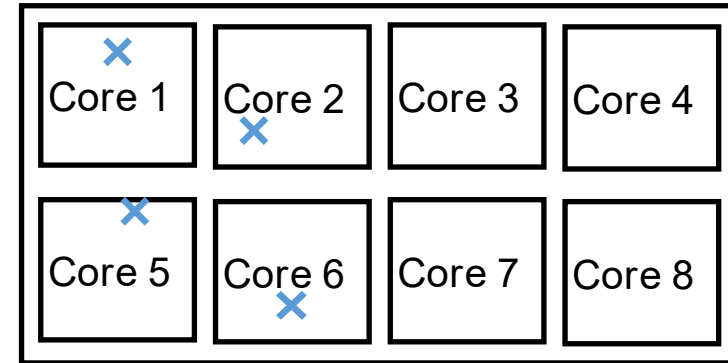
Dependability via Redundancy

- Applies to everything from datacenter to storage to memory to instructors
 - Redundant datacenters so that can lose 1 datacenter but Internet service stays online
 - Redundant disks so that can lose 1 disk but not lose data (Redundant Arrays of Independent Disks/RAID)
 - Redundant memory bits so that can lose some bits but not the data (Error Correcting Code/ECC Memory)



Great Idea 7: Dependability via Redundancy

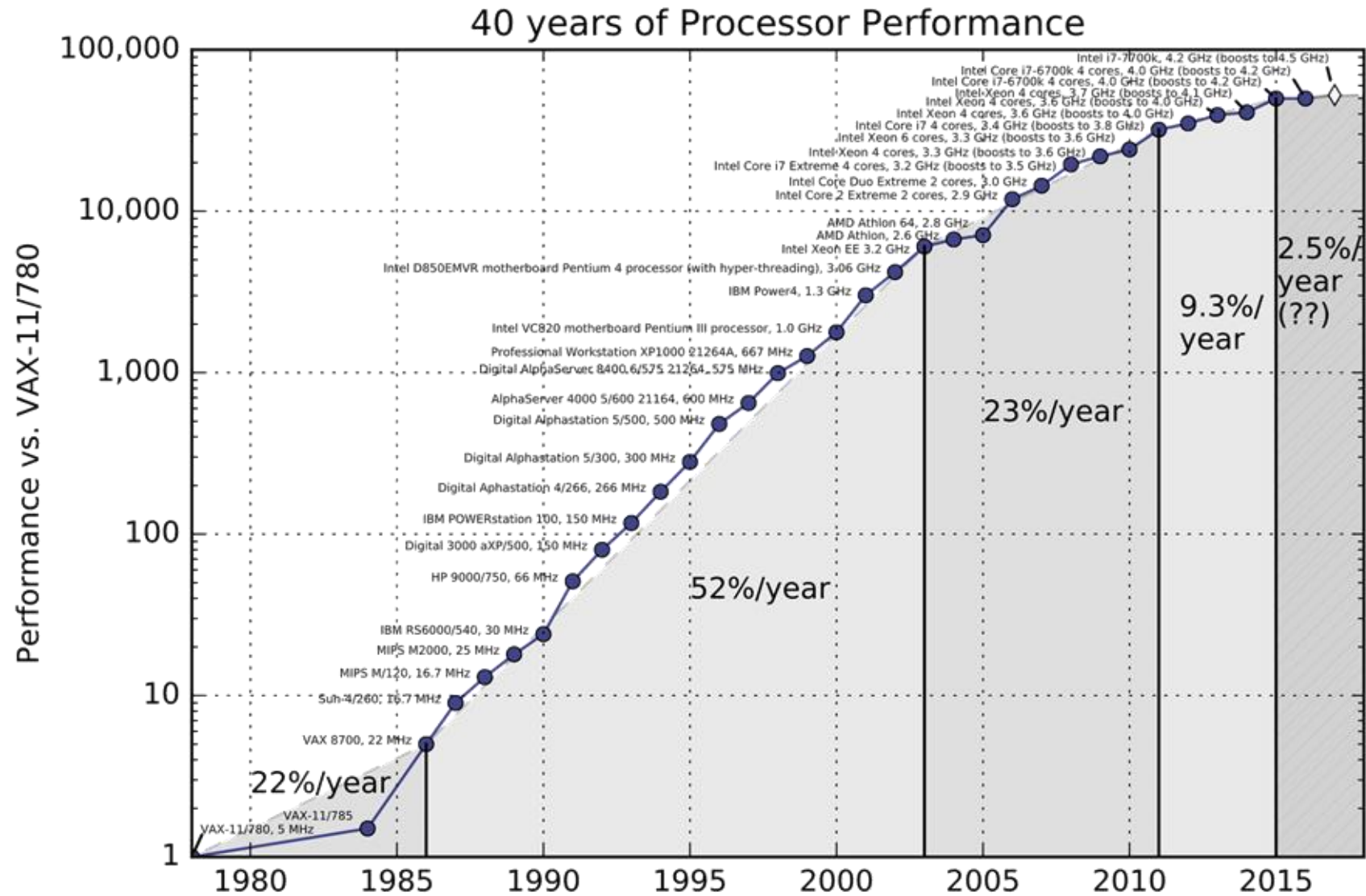
- Redundancy improves yield
 - Considering a CPU with 10^9 transistors
 - Each transistor has a probability 10^{-9} to be failure during production



Old Conventional Wisdom

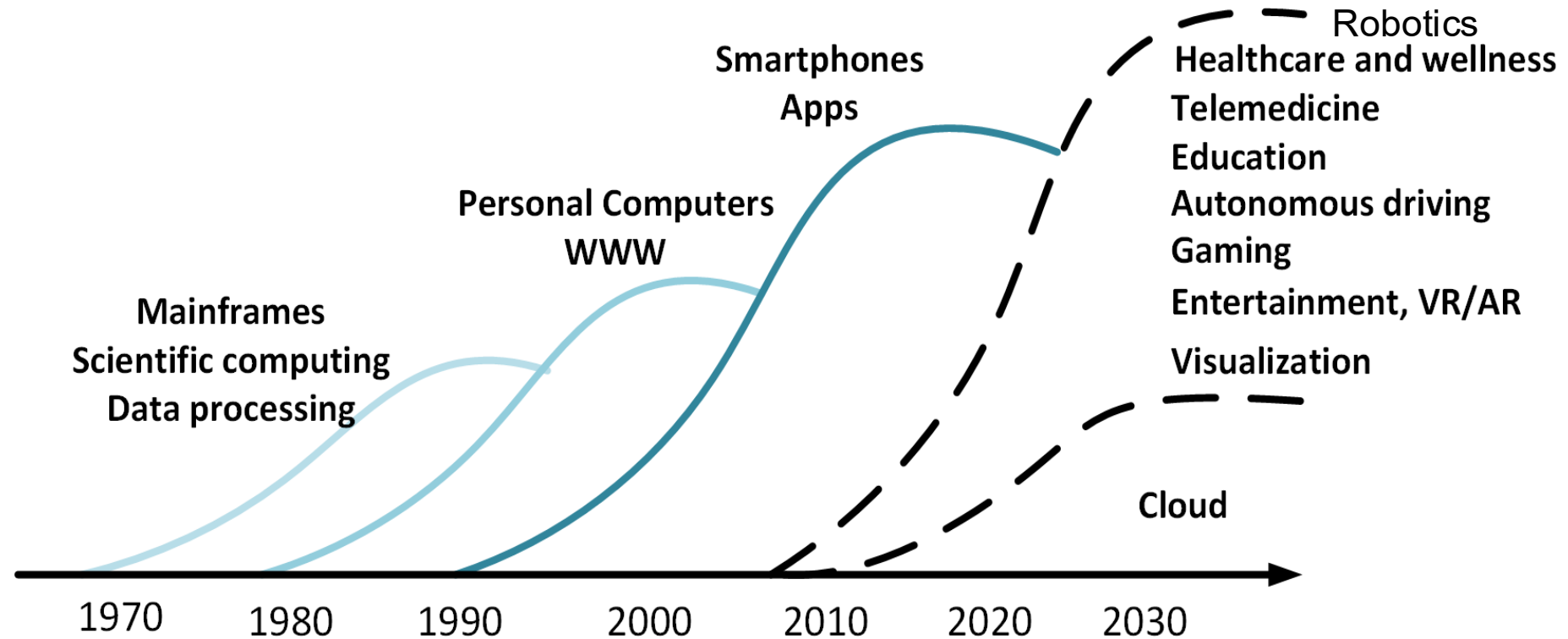
- Moore's Law + Dennard Scaling = faster, cheaper, lower-power general-purpose computers each year
- In glory days, 1%/week performance improvement!
- Dumb to compete by designing specialized computers
- By time you've finished design, next generation of general-purpose will beat you

Why is Architecture Exciting Today?



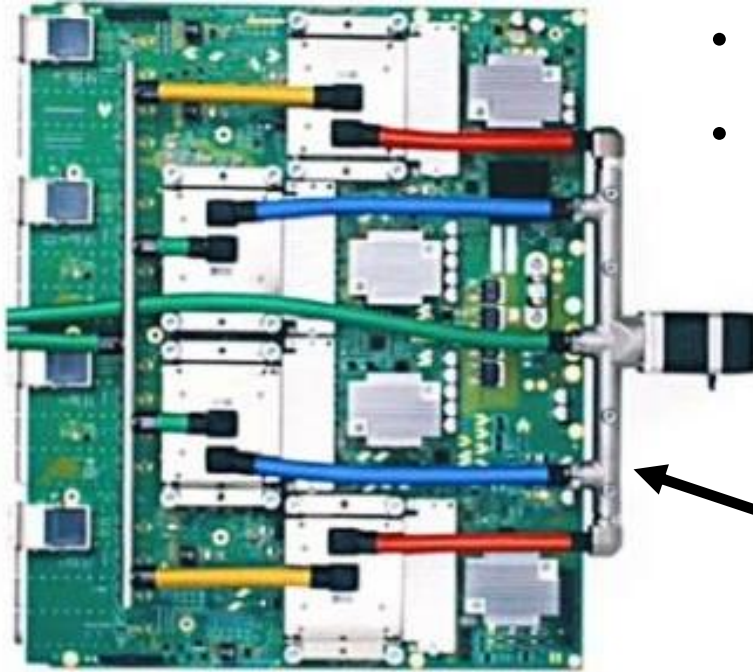
Hennessey & Patterson, 2017

Why is Architecture Exciting Today?



- Number of deployed devices continues growing, but no single killer app
- Diversification of needs, architectures

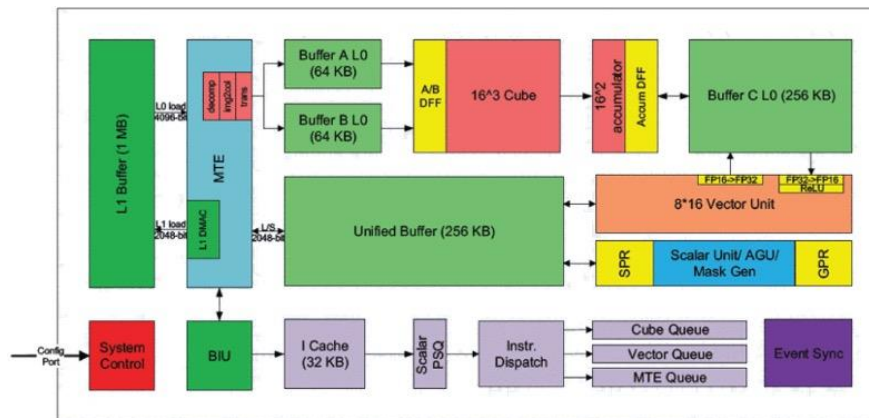
New Conventional Wisdom



- Modern CPUs massively parallel
- Many specialized chips (heterogeneous)

Google TPUv4
Specialized Engine for NN
training
Deployed in cloud
180+ TFLOPS/chip
Water Cooling!

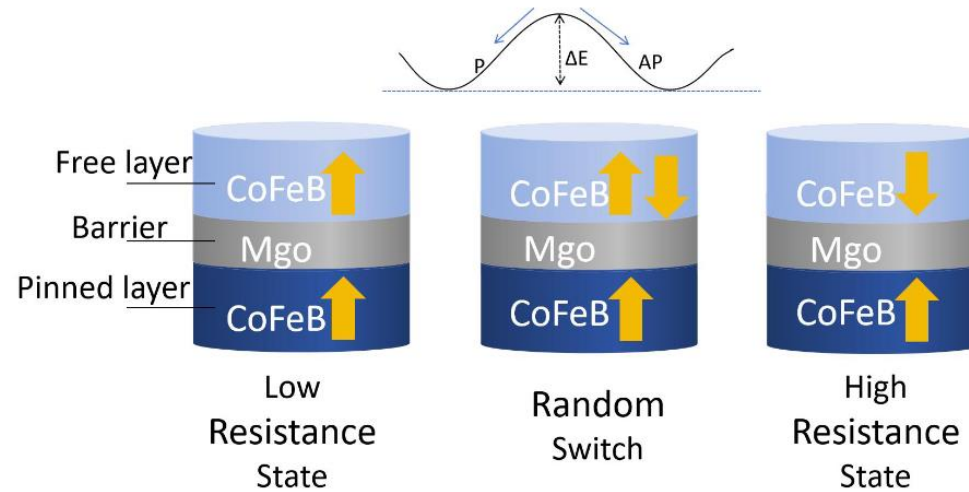
Domain-specific
architectures



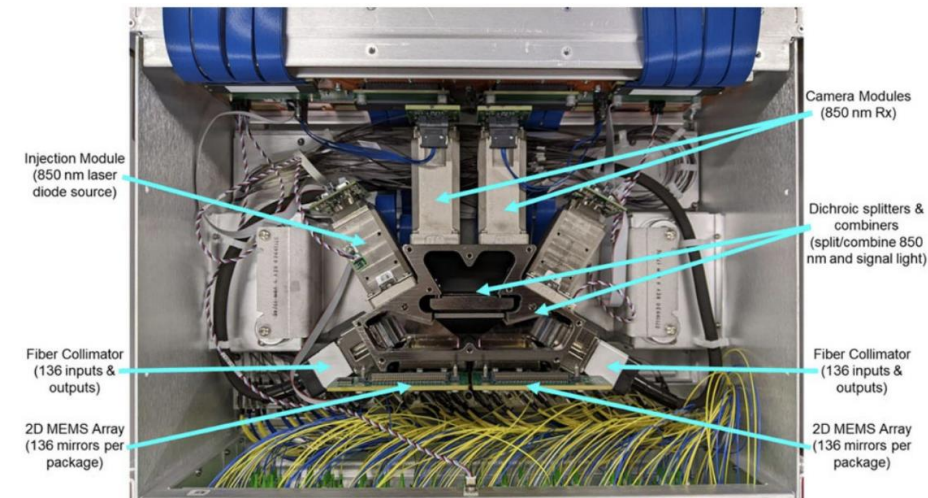
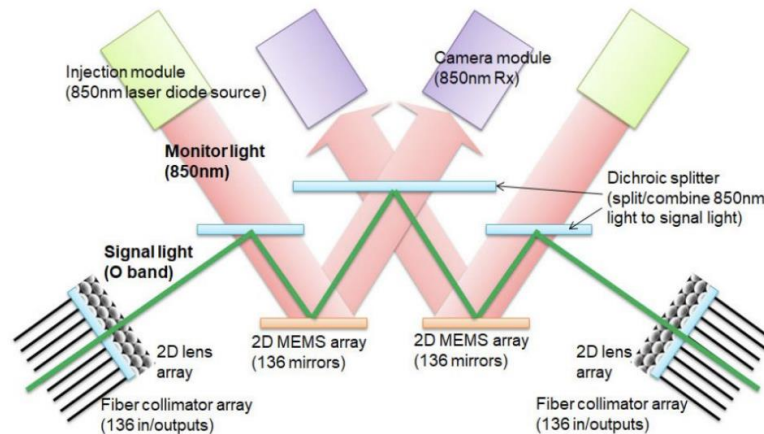
Huawei edge device
Specialized for NN inference
11 TFLOPS FP16; 22 TOPS
INT8

New Conventional Wisdom

- Emerging technologies and devices, e.g.



MTJ (Magnetic Tunnel Junction)
MRAM
Compute-in-memory



Optically Reconfigurable Interconnect inside Google AI infrastructure

Summary

- Great ideas in CA
 - Abstraction (Layers of Representation / Interpretation)
 - Moore's Law (designing through trends)
 - Make the common case fast (Amdahl's Law)
 - Principle of locality (memory hierarchy)
 - Parallelism (pipeline as special case)
 - Performance measurement and improvement
 - Dependability via redundancy