



CENG 5410

Advanced Computer Architecture

Lecture 02: Technology & Performance

Zhengrong Wang
CSE Department, CUHK
zhengrongwang@cuhk.edu.hk



How do we build a high-**performance** computer?

Subject to:

- Cost,
- Power,
- Energy,
- Reliability,
- Security,
- Generality



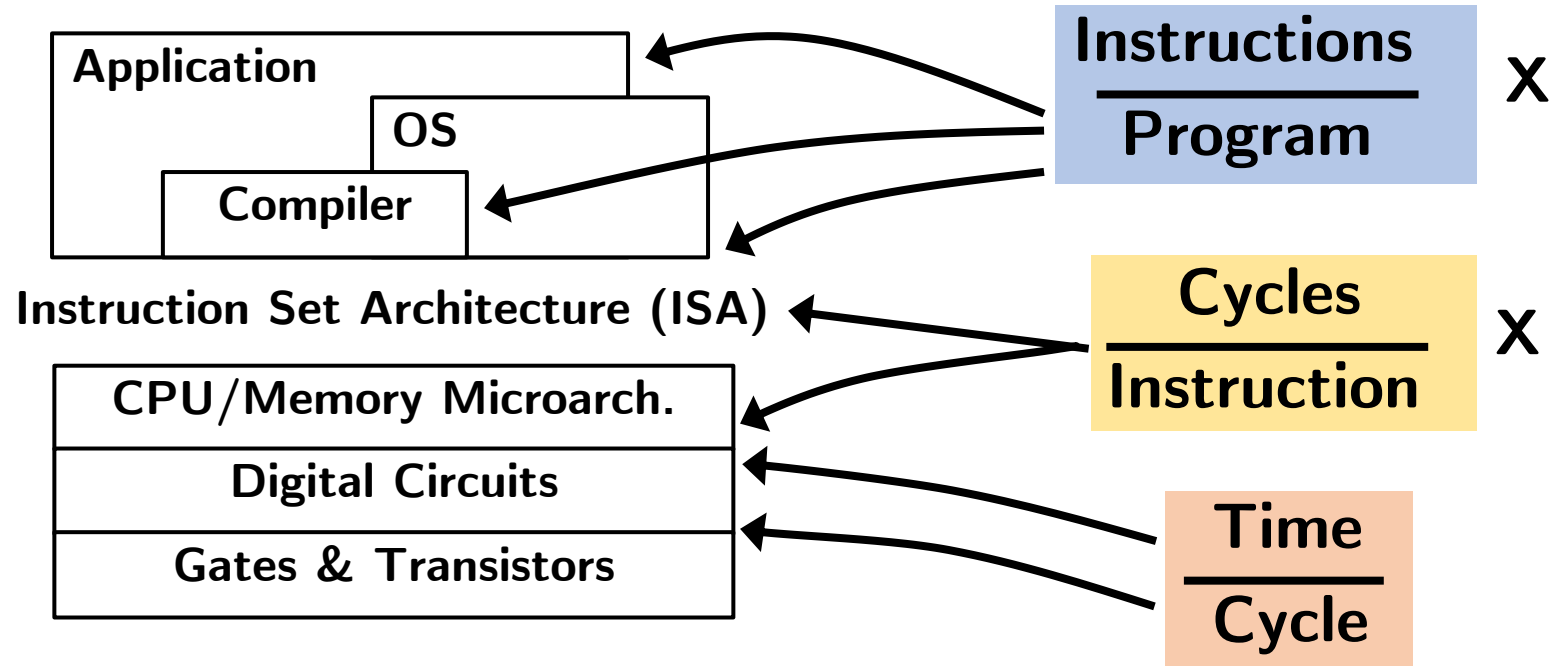
Iron Law of Performance

$$\frac{\text{Time}}{\text{Program}} = \frac{\text{Instructions}}{\text{Program}} \times \frac{\text{Cycles}}{\text{Instruction}} \times \frac{\text{Time}}{\text{Cycle}}$$

- Programs consist of simple operations (instructions)
 - Add two numbers, fetch data value from memory, etc.
- **Instructions per program**: “dynamic instruction count”
 - Runtime count of instructions executed by the program
 - Determined by program, compiler, instruction set architecture (ISA)
- **Cycles per instruction**: “CPI” (typical range: 4 to 1/4)
 - On average, how many *cycles* does an instruction take to execute?
 - Determined by program, compiler, ISA, micro-architecture
- **Seconds per cycle**: clock period, length of each cycle
 - Inverse metric: cycles per second (Hertz) or cycles per ns (Ghz)
 - Determined by micro-architecture, **technology parameters**



Computer System Layers



Main Focus for Today

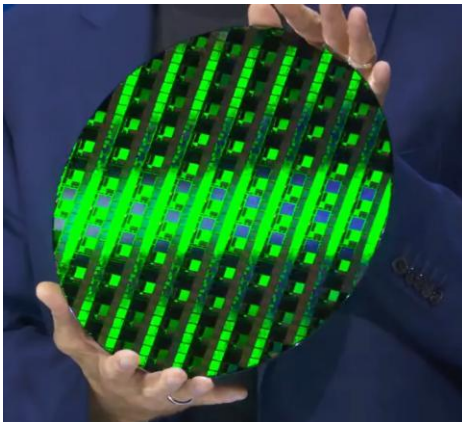


Computer Building Blocks

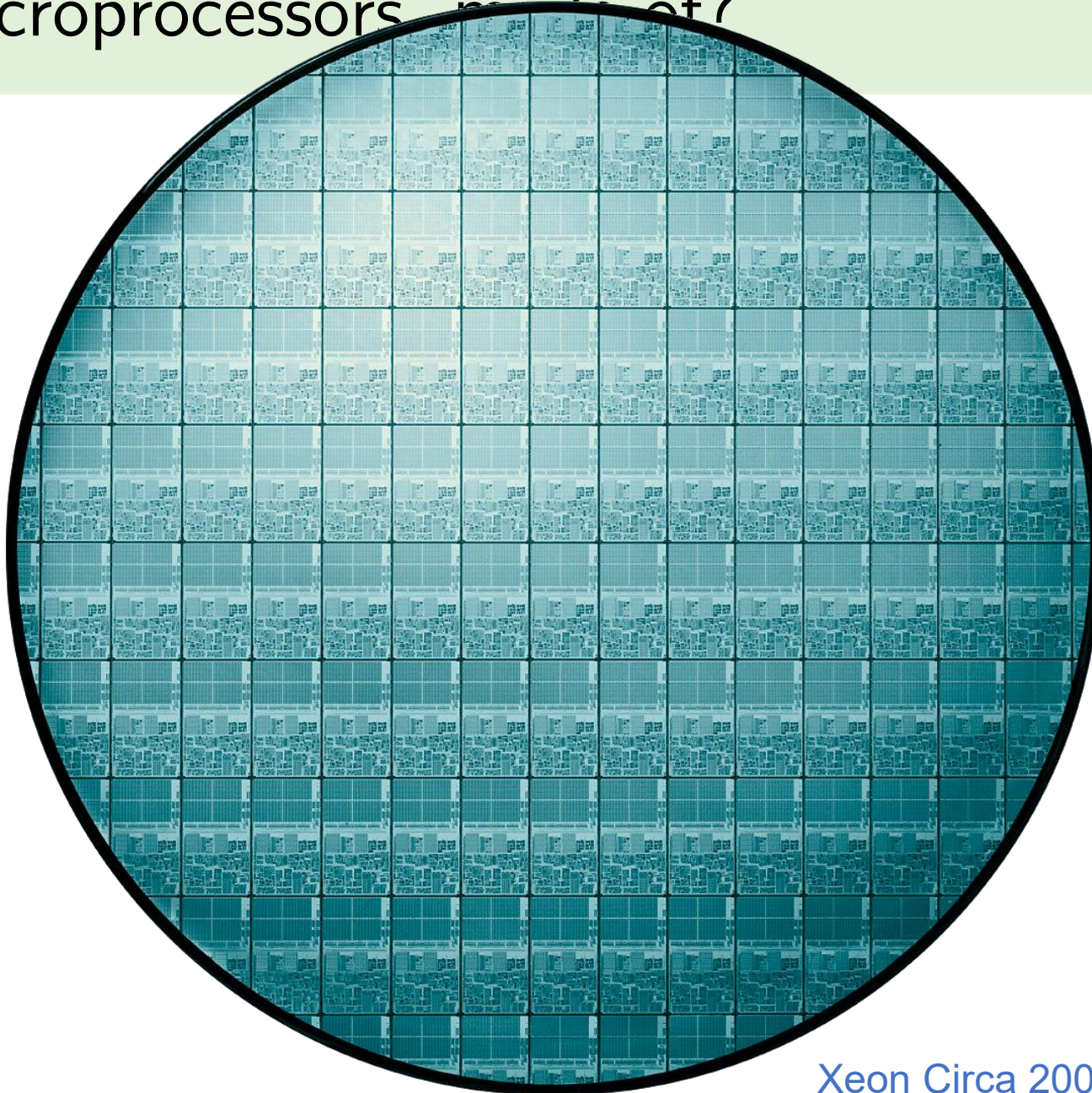
- Transistors/Gates (computing)
 - How can they be connected to do something useful?
 - How do we evaluate how fast a logic block is?
- Wires (transporting)
 - What and where are they?
 - How can they be modeled?
- Memories (storing)
 - SRAM (speed) vs. DRAM (density)
- Compute and Store:
 - CAMs: Memories enabling associative lookups (lookup by value rather than by index)
 - Processing-in-memory
 - Digital: vectorized logical ops over some (2-3) RAM bitlines
 - Analog: arithmetic ops over many RAM bitlines
 - e.g. Memristors: Memories with programmable resistance



What are microprocessors made of?



18a wafer, 2024

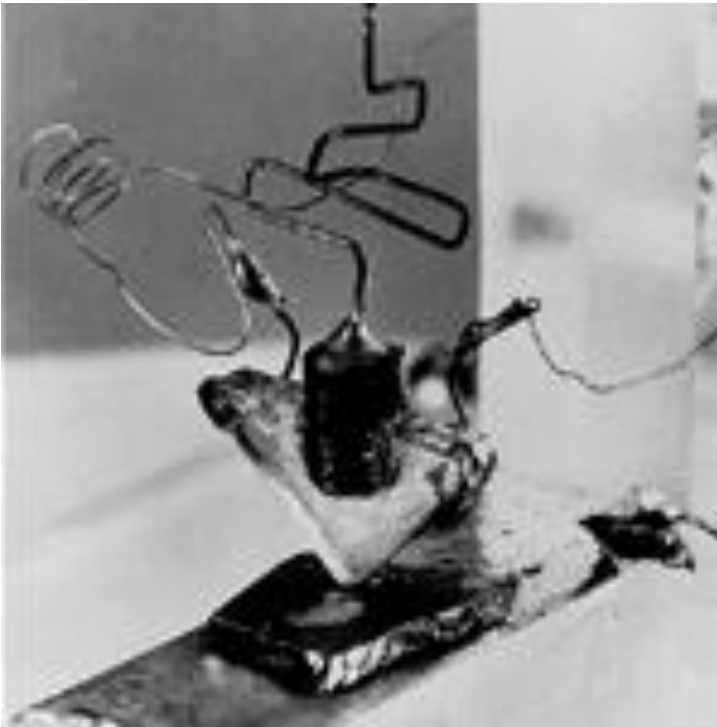


Xeon Circa 2005

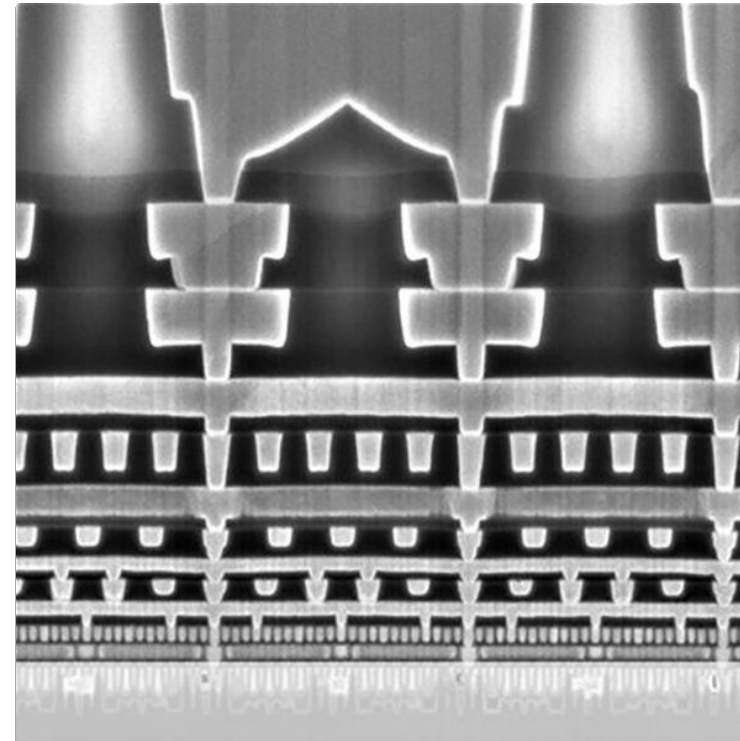


Semiconductor Technology Background

Transistor Circa 1947



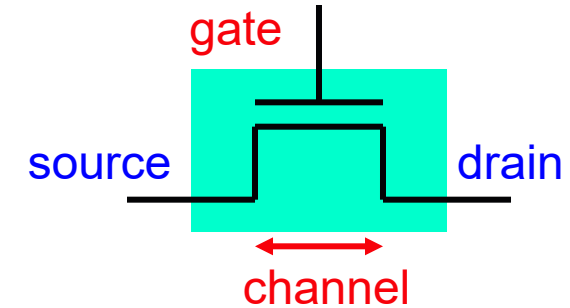
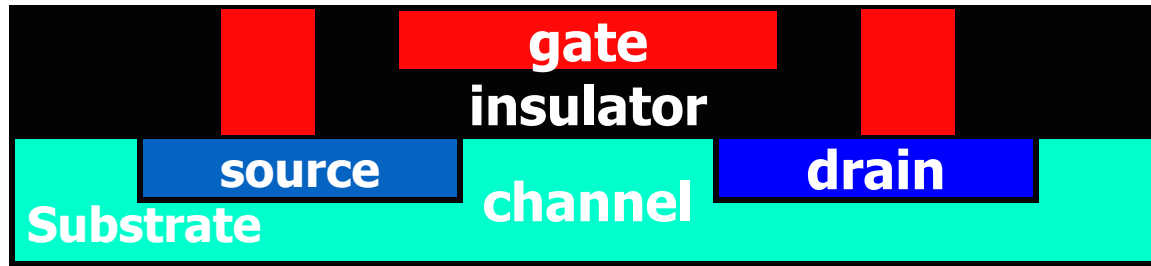
Transistor Circa Now



Human hair $\sim 100\mu\text{m}$,
(50000 bigger than 2nm transistor)



Basic technology element: MOSFET



- Solid-state component acts like electrical switch
- **MOS**: three ingredients for the transistor
 - **Metal**: Aluminum, Tungsten, Titanium Nitride: **conductor**
 - **Oxide**: Originally Silicon Dioxide: **insulator**
 - **Semiconductor**: doped Si: conducts under certain conditions
- **FET**: field-effect (the mechanism) transistor
 - Voltage on gate: current flows source to drain (transistor on)
 - No voltage on gate: no current (transistor off)
- **Tech node (length)**: characteristic parameter (short → fast)
 - Aka “feature size” or “technology”
 - Currently: ~2 nm (but doesn’t really measure anything anymore)



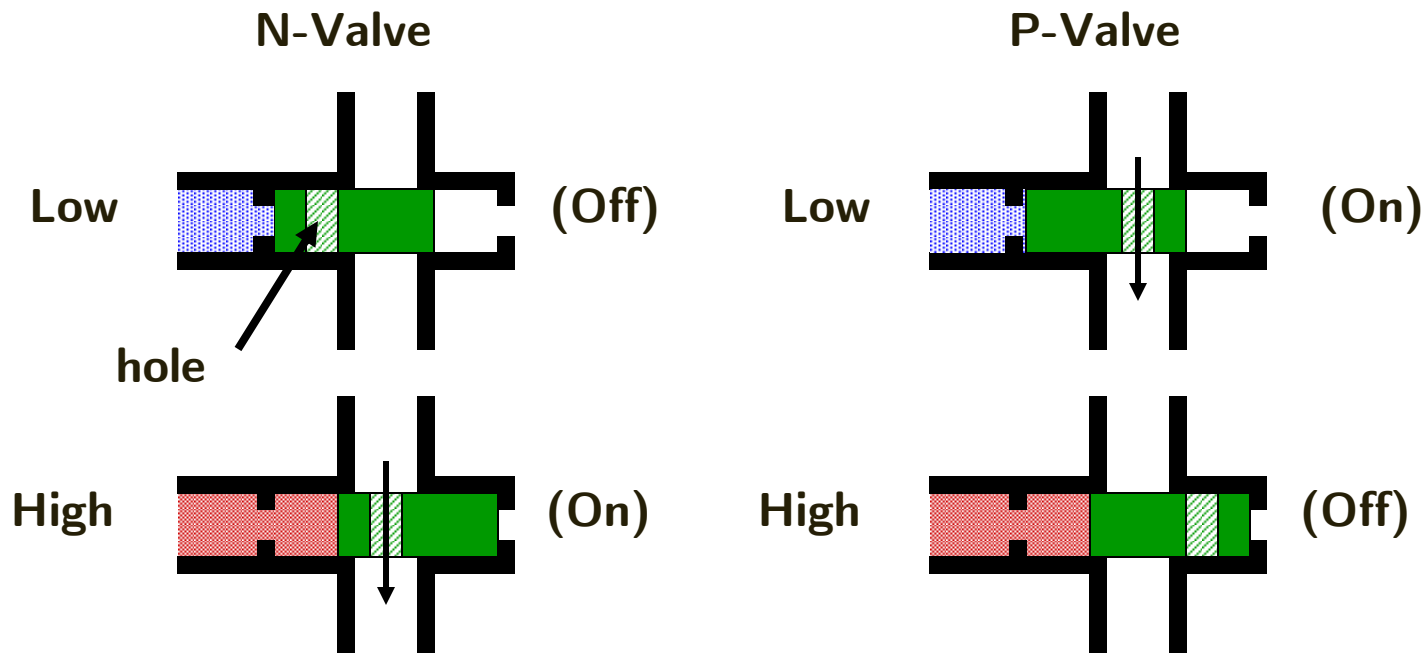
How do Transistors work?

An analogy for CS simpletons (i.e. me :)



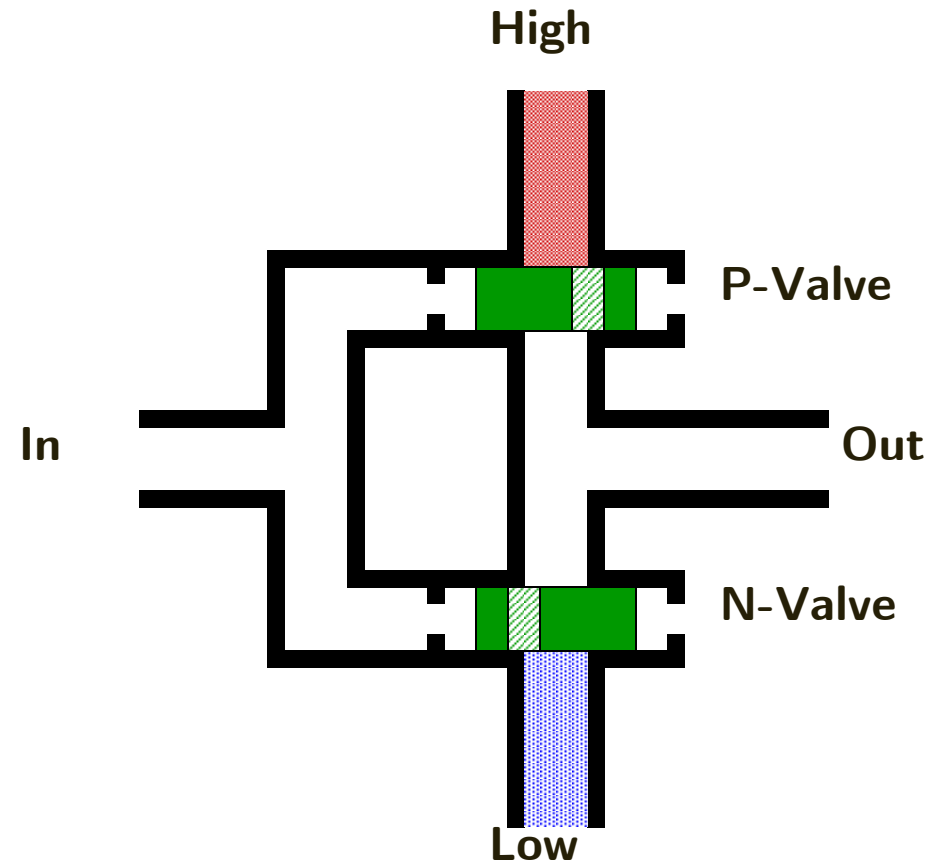
A Transistor Analogy: Computing with Air

- Use air pressure to encode values
 - High pressure represents a “1” (blow)
 - Low pressure represents a “0” (suck)
- Valve can allow or disallow the flow of air
 - Two types of valves





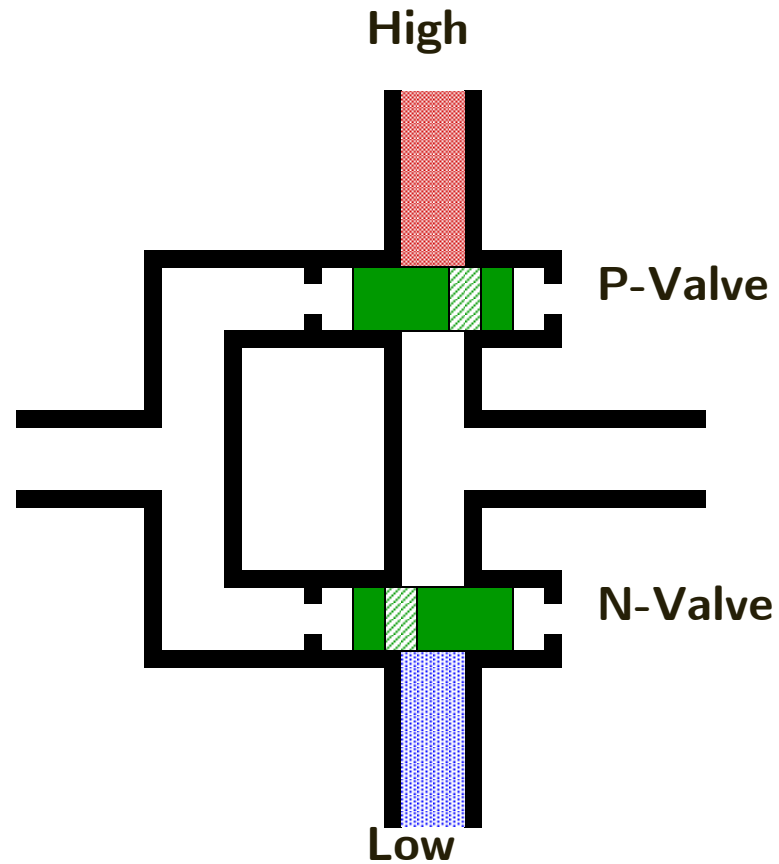
Pressure Inverter





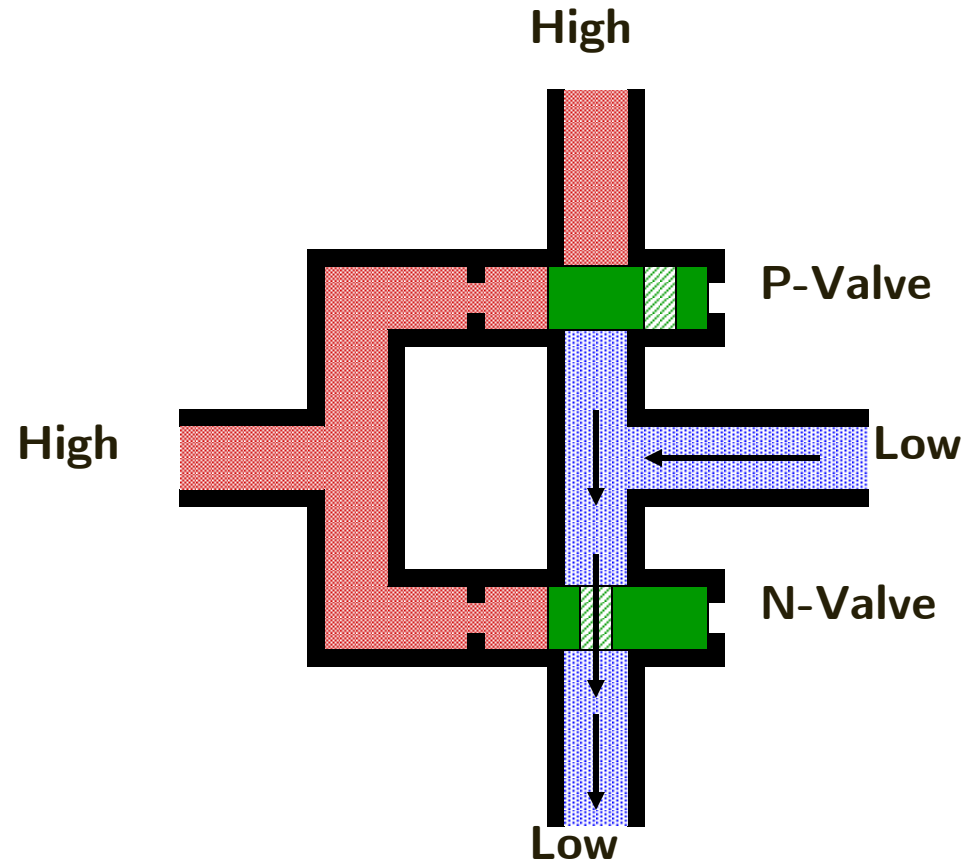


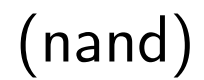
Pressure Inverter





Pressure Inverter (High to Low)







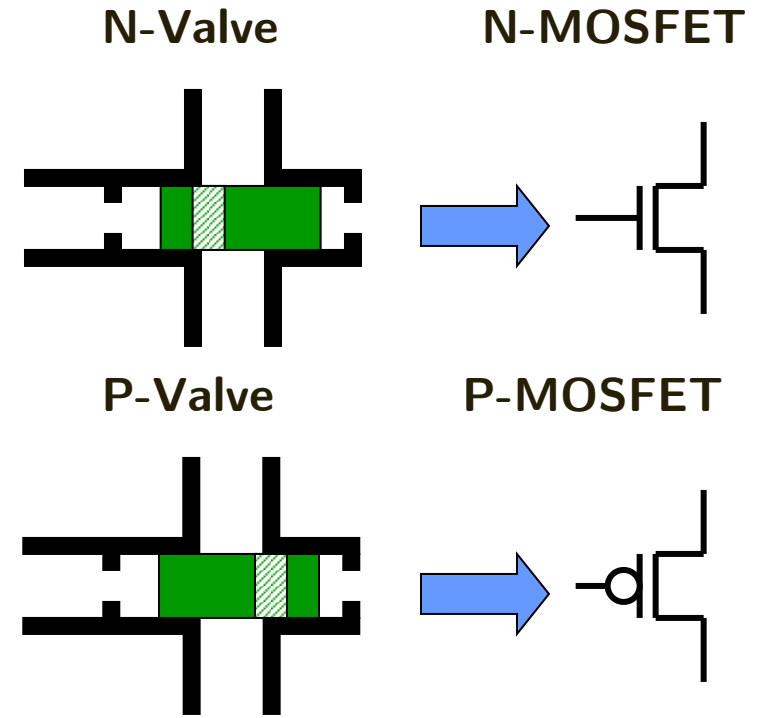
Analogy Explained

- Pressure differential → electrical potential (voltage)
 - Air molecules → electrons
 - Pressure (molecules per volume) → voltage
- Air flow → electrical current (eh, good enough)
 - Pipes → wires
 - Air/Electrons only flows from high to low pressure/voltage
 - Flow only occurs when changing from 1 to 0 or 0 to 1
 - Lesson: Transitions costs energy!
- Valve → transistor
- The transistor: one of the century's most important inventions



Transistors as Switches

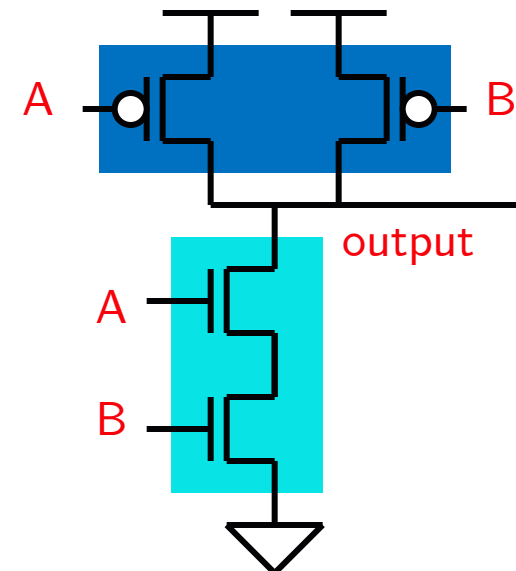
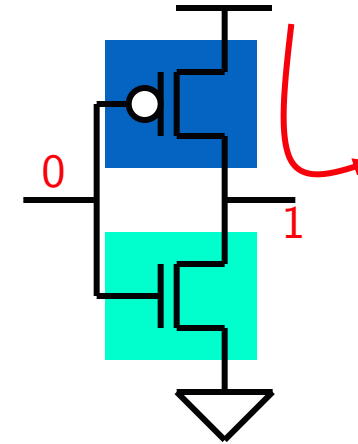
- Two types
 - N-type: Conduct when gate voltage is 1
 - P-type: Conduct when gate voltage is 0
- Properties
 - Solid state (no moving parts)
 - Reliable (low failure rate)
 - Small (2nm channel length)
 - Fast ($<0.1\text{ns}$ switch latency)
- **CMOS**: complementary n-/p-networks form boolean logic





CMOS Examples

- Example I: inverter
 - Case I: input = 0
 - P-transistor closed, n-transistor open
 - Power charges output (1)
 - Case II: input = 1
 - P-transistor open, n-transistor closed
 - Output discharges to ground (0)
- Example II: look at **truth table**
 - $0, 0 \rightarrow 1$ $0, 1 \rightarrow 1$
 - $1, 0 \rightarrow 1$ $1, 1 \rightarrow 0$
 - Result: this is a **NAND** (NOT AND)
 - NAND is universal (can build any logic function)

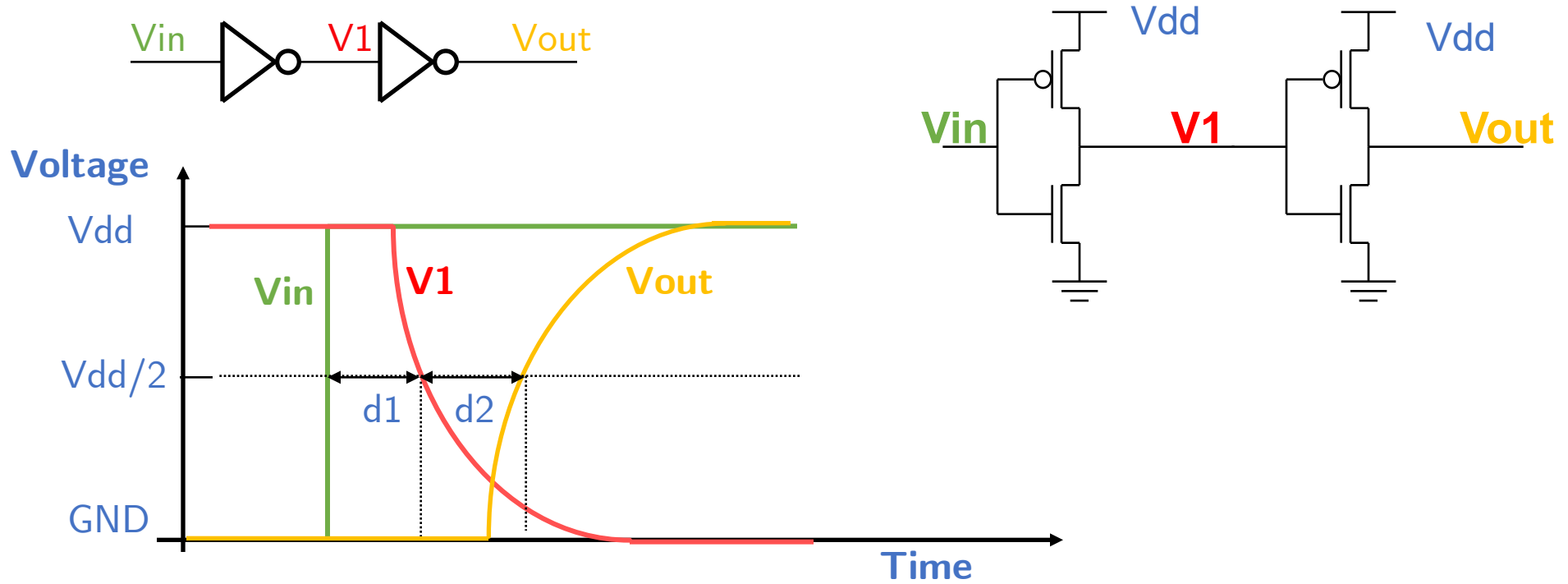




Circuit Timing & Critical Path



What about Transistor Delay

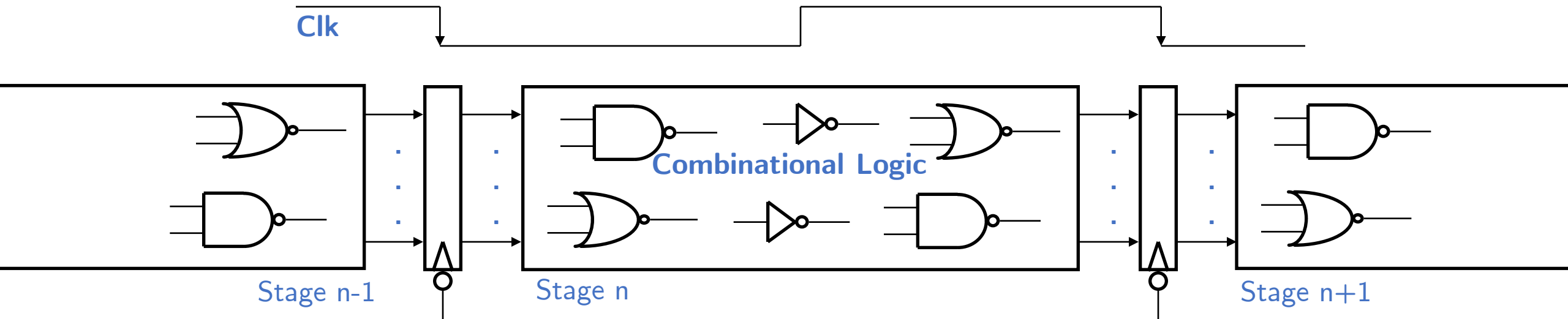


- Total Propagation Delay = Sum of individual delays = $d_1 + d_2$



Synchronous Design

- Hard to coordinate design components with logic only
 - “asynchronous” design – rarely used in CPUs/GPUs
- Synchronous: State elements synchronized by clock

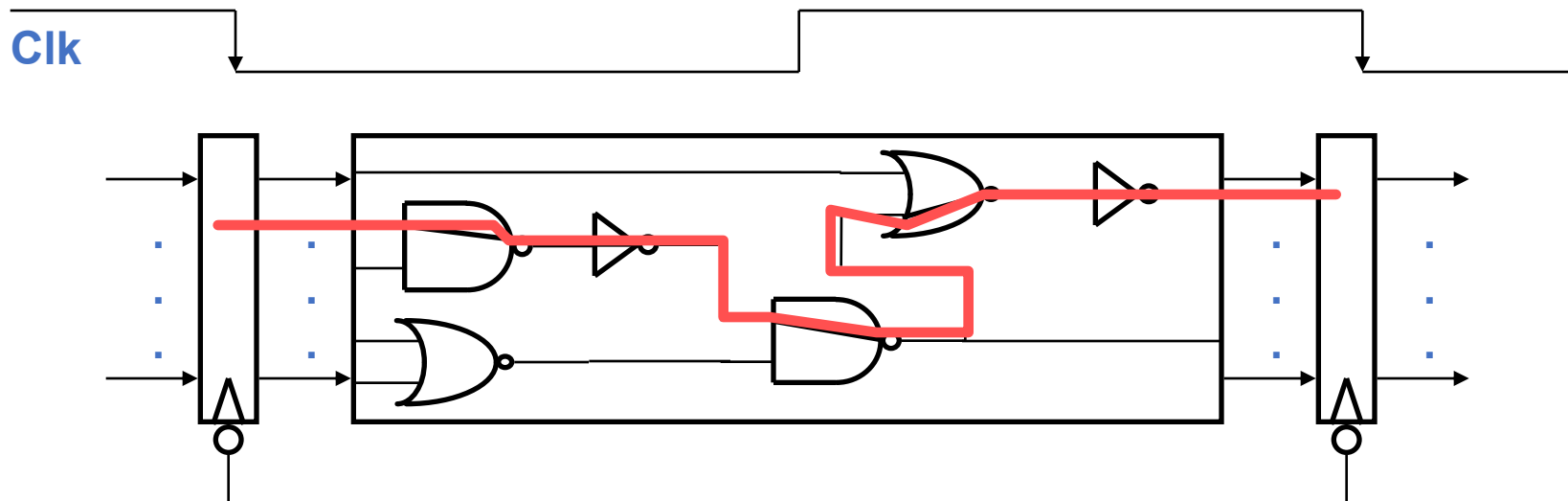


- All storage elements are clocked by the same clock edge
- The combination logic block's:
 - Inputs change after a clock edge and propagate through
 - All outputs **MUST** be stable before the next clock tick



Critical Path & Cycle Time

- Circuit Critical path: the slowest combinational path between any two storage devices
- Cycle time (aka. clock period) $>$ circuit critical path time
- **This is where frequency comes from :)**
 - $\text{Frequency} = 1 / \text{clock_period}$





Technology Basis of Transistor Speed

- Physics 101: delay through an electrical component $\sim RC$

- **Resistance (R)**

- Slows rate of charge flow



$\sim$ length / cross-section area

- **Capacitance (C)**

- Stores charge



$\sim$ increases w/ area, length, closer spacing

- **Voltage (V)**

- Electrical pressure

- **Threshold Voltage (V_t)**

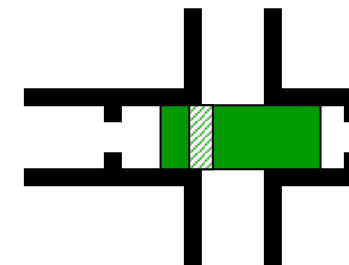
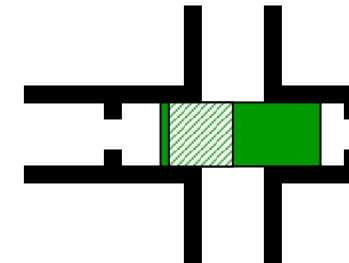
- Voltage at which a transistor turns “on”
 - Property of transistor based on fabrication technology

- **Switching time $\sim$ to $(R * C) / (V - V_t)$**
(first order intuition...)

- Two kinds of electrical components

- CMOS transistors (gates)
 - Wires

Low V_t (faster, more power)

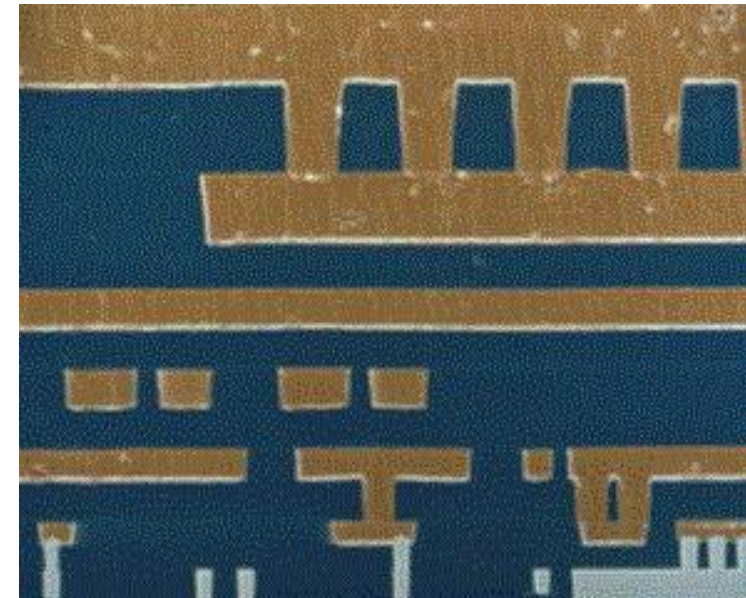
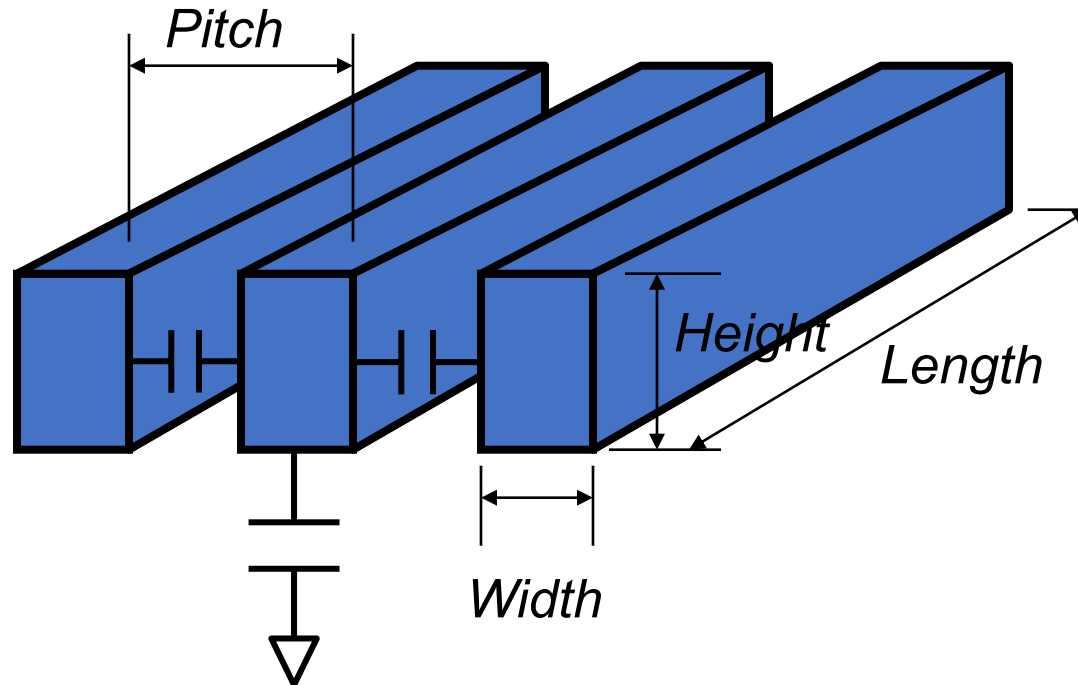


High V_t (slower, less power)



Wire Geometry

- Wires 4-dimensional: **length**, **width**, **height**, **“pitch”**
 - Longer wires have more resistance
 - “Thinner” wires have more resistance
 - Closer wire spacing (“pitch”) increases capacitance



IBM CMOS7, 6 layers of copper wiring



Increasing Problem: Wire Delay

- RC Delay of wires
 - **Resistance** proportional to: $\text{resistivity} * \text{length} / (\text{cross section})$
 - Wires with smaller cross section have higher resistance
 - Resistivity (type of metal, copper vs aluminum)
 - **Capacitance** proportional to length
 - And wire spacing (closer wires have large capacitance)
 - Permittivity or “dielectric constant” (of material between wires)
- Result: delay of a wire is **quadratic** in length
 - Insert “inverter” repeaters for long wires
 - Why? To bring it back to linear delay... but repeaters still add delay
- Trend 1: wires are getting relatively slow to transistors
- Trend 2: chips are getting larger
- **Arch. Implication:**
 - **Avoid limit the amount of communication**
 - **Keep high-bandwidth communication as local as possible**



Power & Energy



Power/Energy Are Increasingly Important

- **Battery life** for mobile devices
 - Laptops, phones, cameras
- **Tolerable temperature** for devices without active cooling
 - Power means temperature, active cooling means **cost**
 - No room for a fan in a cell phone, no market for a hot cell phone
- **Electric bill** for compute/data centers
 - Pay for power twice: once in, once out (to cool)
- **Environmental concerns**
 - “Computers” account for growing fraction of energy consumption
 - As a percentage of global carbon emissions
 - Bitcoin: ~0.1%
 - Smartphones: ~0.3%
 - Datacenters: ~3% (caveat: a lot of it is manufacturing)
 - Commercial Flights: ~2.4%



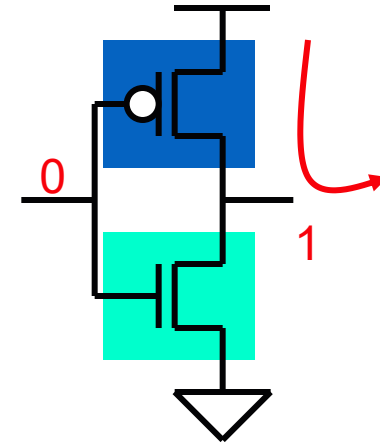
Energy & Power

- **Energy**: measured in Joules or Watt-seconds
 - Total amount of energy stored/used
 - Battery life, electric bill, environmental impact
 - Joules per Instruction, or Joules per Operation (car analogy: gallons per mile)
- **Power**: energy per unit time (measured in Watts)
 - Joules per second (car analogy: gallons per hour)
 - Power impacts power supply and cooling requirements (cost)
 - Power-density (Watt/mm^2): important related metric
 - Peak power vs average power
 - E.g., camera, power “spikes” when you actually take a picture
- Two sources:
 - **Dynamic power**: active switching of transistors
 - **Static power**: leakage of transistors even while inactive



Dynamic Power

- **Dynamic power (P_{dynamic}):** aka switching or active power
 - Energy to switch on a gate (0 to 1)
 - Each gate has capacitance (C)
 - Energy to charge/discharge a capacitor is $\sim$ to $C * V^2$
- **$P_{\text{dynamic}} \sim N * C * V^2 * f * A$**
 - N: number of transistors
 - C: capacitance per transistor (size of transistors)
 - V: voltage (supply voltage for gate)
 - f: frequency (transistor switching freq. is $\sim$ to clock freq.)
 - A: activity factor (not all transistors may switch this cycle)





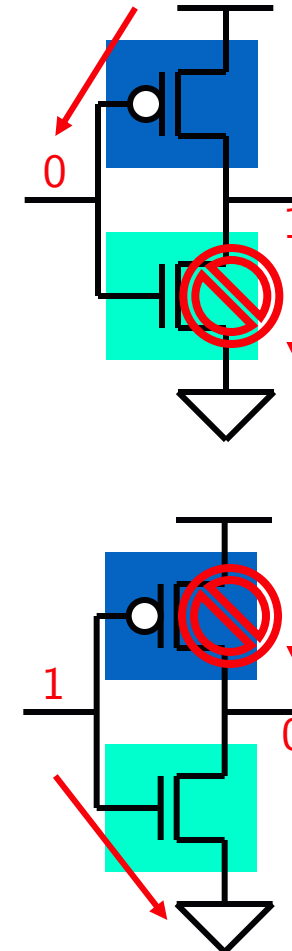
Reducing Dynamic Power

- Target each component: $P_{\text{dynamic}} \sim N * C * V^2 * f * A$
- **Reduce number of transistors** (N)
 - Use fewer transistors/gates (better design; specialized hardware)
- **Reduce capacitance** (C)
 - Smaller transistors (Moore's law)
- **Reduce voltage** (V)
 - Quadratic reduction in energy consumption!
 - But also slows transistors (recall transistor speed is $\sim$ to V)
- **Reduce frequency** (f)
 - Slower clock frequency (reduces power but not energy) Why?
- **Reduce activity** (A)
 - "Clock gating" disable clocks to unused parts of chip
 - Don't switch gates unnecessarily



Static Power

- **Static power (P_{static})**: aka idle or leakage power
 - Transistors don't turn off all the way
 - Transistors "leak"
 - Analogy: leaky valve
 - $P_{\text{static}} \sim N * V * e^{-V_t}$
(don't take this too literally)
 - N: number of transistors
 - V: voltage
 - **V_t (threshold voltage)**: voltage at which transistor conducts (begins to switch)
- Switching speed vs leakage trade-off
- The lower the **V_t** :
 - Faster transistors (rough linear) $\sim V - V_t$
 - Leakier transistors (exponential)





Reducing Static Power

- Target each component: $P_{\text{static}} \sim N * V * e^{-V_t}$
- **Reduce number of transistors** (N)
 - Use fewer transistors/gates
- **Disable transistors** (also targets N)
 - “Power gating” disable power to unused parts (long latency to power up)
 - Power down units (or entire cores) not being used
- **Reduce voltage** (V)
 - Linear reduction in static energy consumption
 - But also slows transistors (transistor speed is $\sim$ to V)
- **Dual V_t** – use a mixture of high and low V_t transistors
 - Use slow, low-leak transistors in SRAM arrays
 - Requires extra fabrication steps (cost)
- Note: reducing frequency can actually hurt static energy. Why?



Scale Voltage & Frequency Together

- Recall: $P_{\text{dynamic}} \sim N * C * V^2 * f * A$
- Reduce both voltage and frequency linearly
 - **Cubic decrease in dynamic power**
 - Linear decrease in performance (sub-linear if memory bound)
 - Thus, only about quadratic in energy
 - Linear decrease in static power
 - Thus, only modest static energy improvement
- Example: Intel Xscale (original DVFS experiments)
 - 1 GHz $\rightarrow$ 200 MHz reduces energy used by 30x
 - But around 5x slower
 - 5 x 200 MHz in parallel, use **1/6th the energy**
- **Power-limited designs favor multi-cores!**



Dynamic Voltage/Frequency Scaling (DVFS)

- **Dynamically trade-off performance for lower power**
 - Change the voltage and frequency at runtime
 - Under control of operating system
 - Best of both worlds?
- Modern chips can adjust frequency on a per-core basis

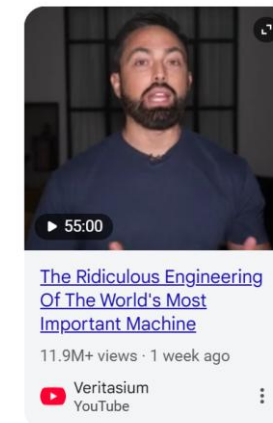
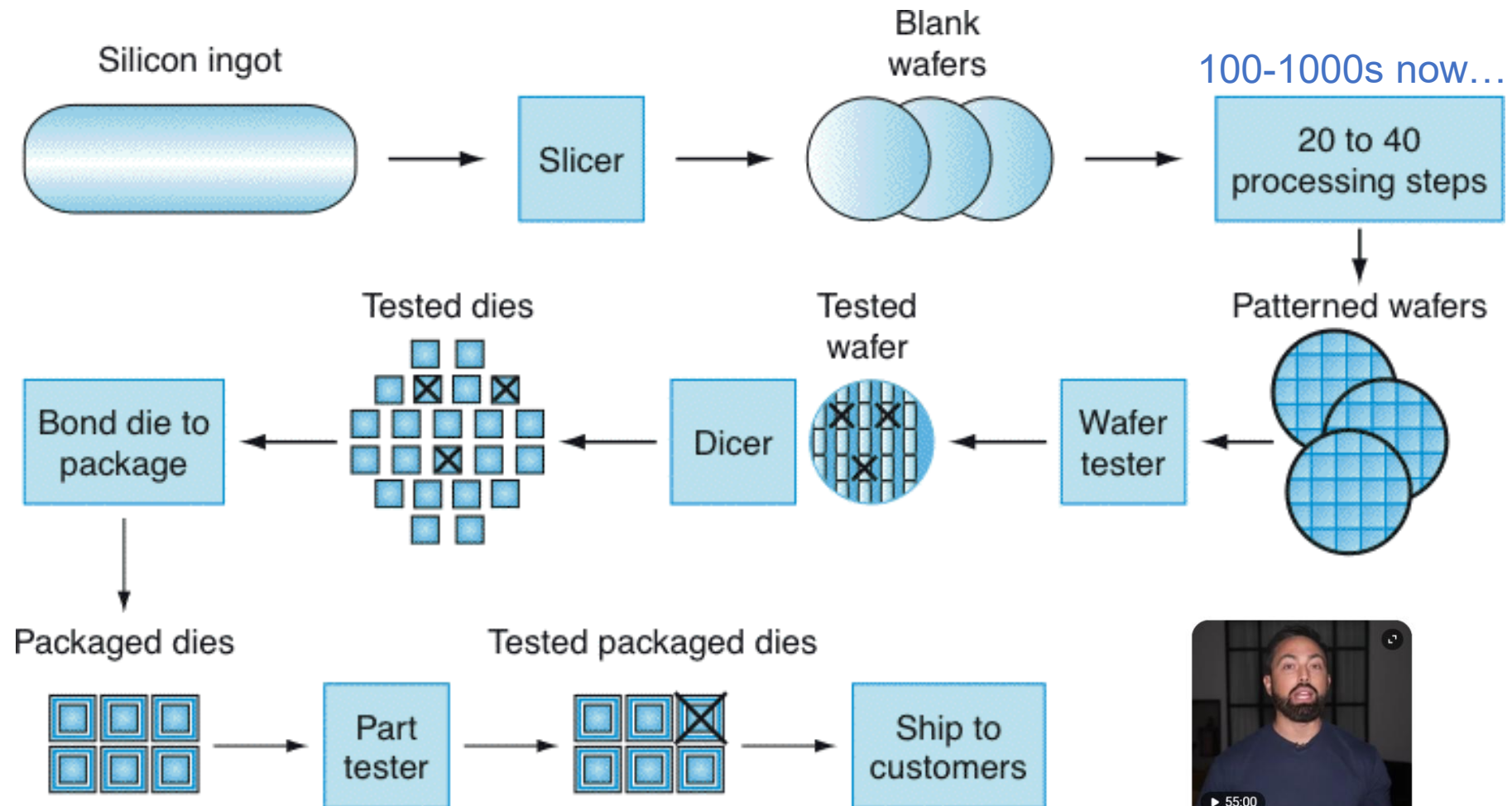
	Mobile PentiumIII	Alder Lake
f (MHz)	300–1000 (step=50)	1000-5200 (step=100)
V (V)	0.9–1.7 (step=0.1)	1.2–1.45V (cont)
Power (W)	~4.5-34W	~30-250W



Manufacturing:
Get those transistors on a chip!

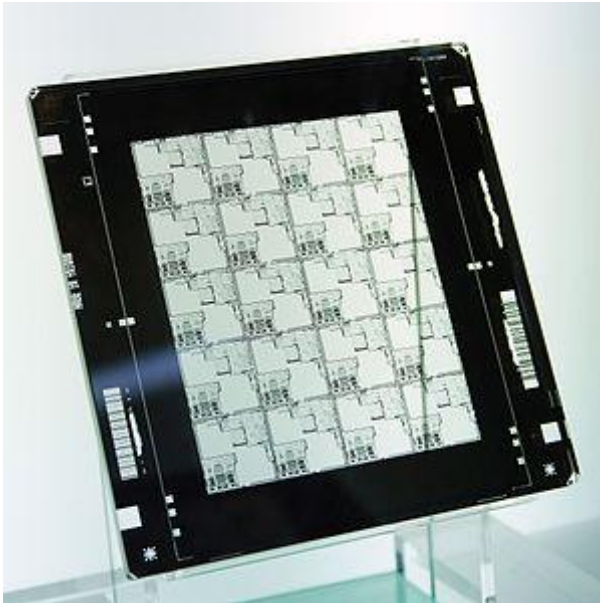


Manufacturing Steps



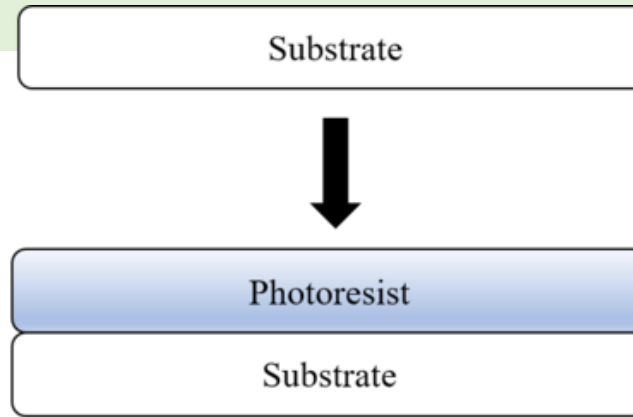


Lithography Step

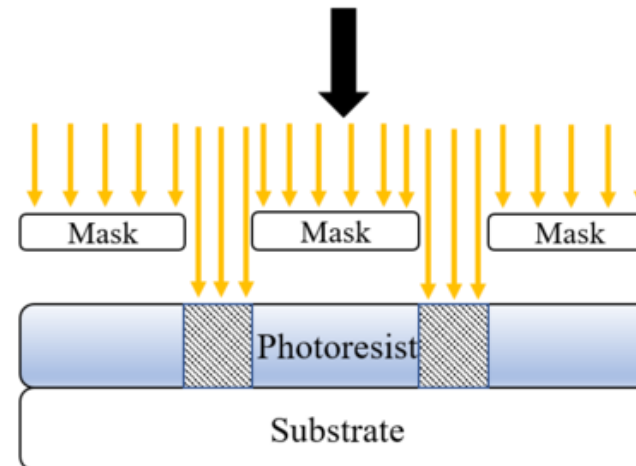


Photomask

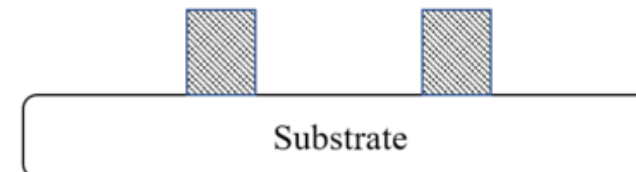
(1) Apply photoresist



(2) Expose to light



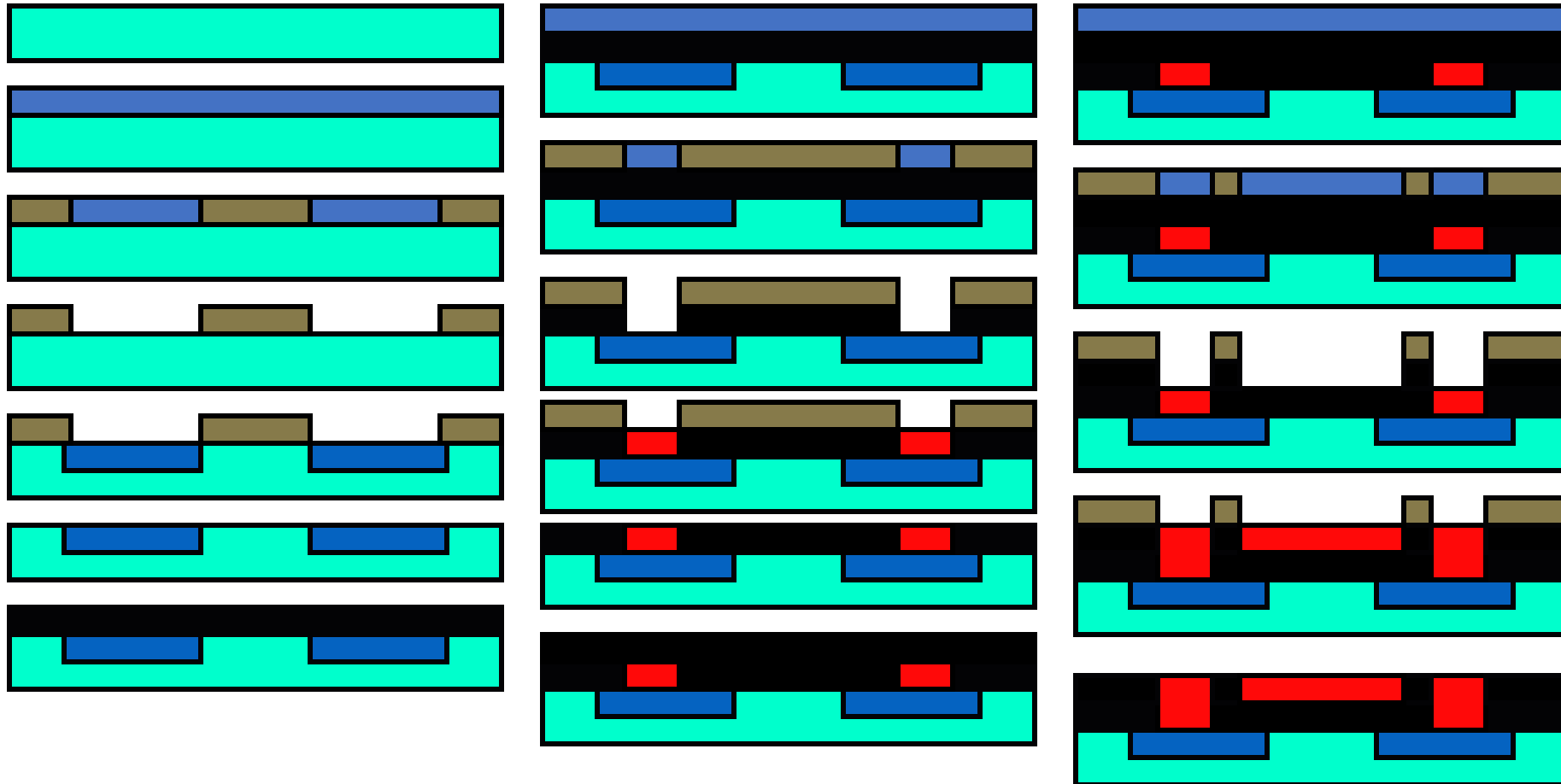
Negative Photoresist





Manufacturing Steps

- Multi-step photo-/electro-chemical process
 - More steps, higher unit cost
- + Fixed cost mass production (\$1 million+ for “mask set”)





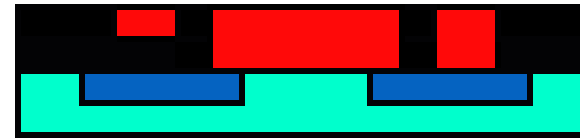
Manufacturing Defects

- Defects can arise
 - Under-/over-doping
 - Over-/under-dissolved insulator
 - Mask mis-alignment
 - Particle contaminants
- Try to minimize defects
 - Process margins
 - Design rules
 - Minimal transistor size, separation
- Or, tolerate defects
 - Redundant or “spare” memory cells
 - Can substantially improve yield

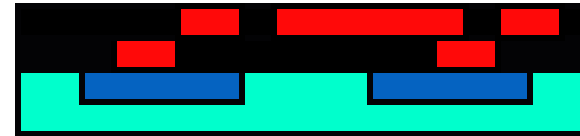
Correct:



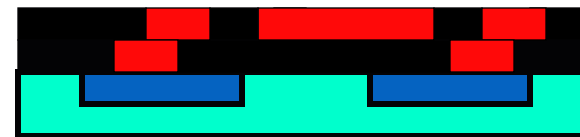
Defective:



Defective:



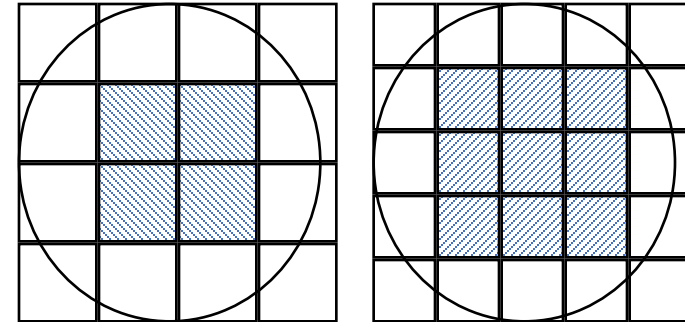
Slow:



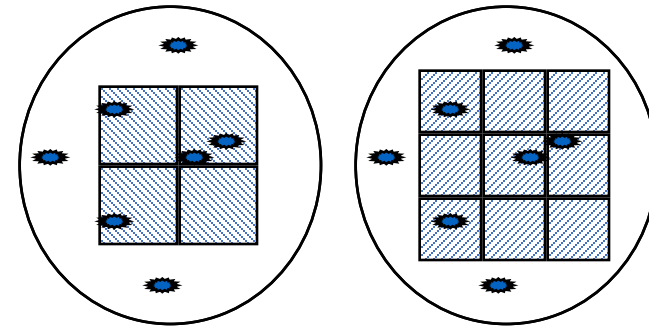


Cost Implications of Defects

- Cost of manufacturing process depends on:
 - Size of the wafer
 - Number of steps
- Cost/mm² of a **chip** depends on *chip area*
 - Why? random defects
 - Larger chip: more chance of defect



- **Wafer yield**: % wafer that is chips
- **Die yield**: % chips that work
- Yield is increasingly non-binary – fast vs slow chips (instead of working vs broken)





All Roads Lead To Multi-Core

+Multi-cores reduce unit costs

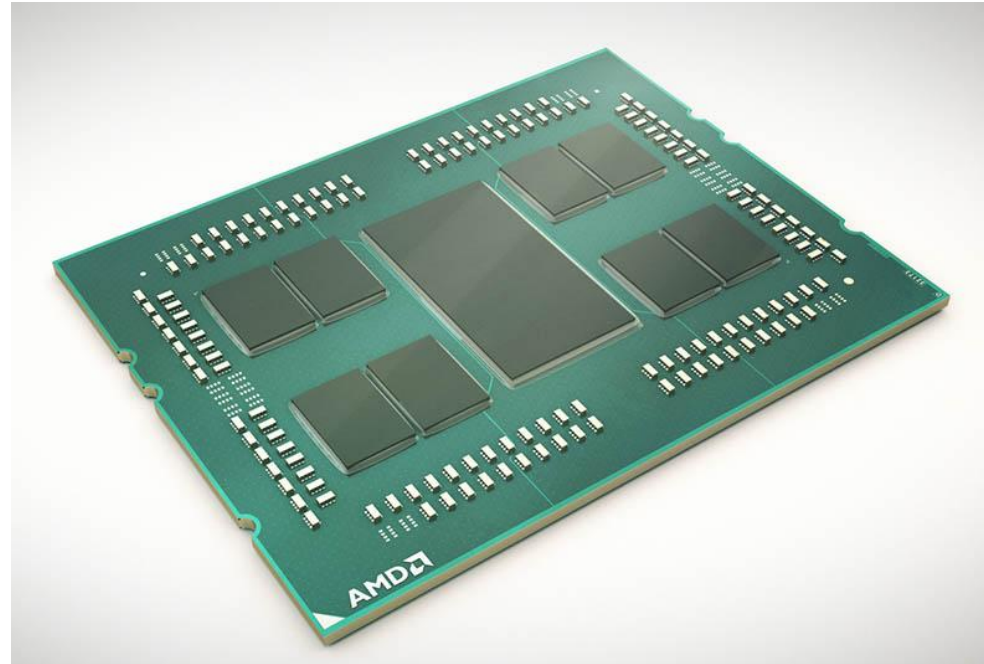
- Higher yield than same-area single-cores
- Why?
 - Defect on one of the cores? Sell remaining cores for less
 - This is part of why you see so many versions of CPUs or GPUs in a single generation with slightly different core counts

+Multi-cores can reduce design costs too

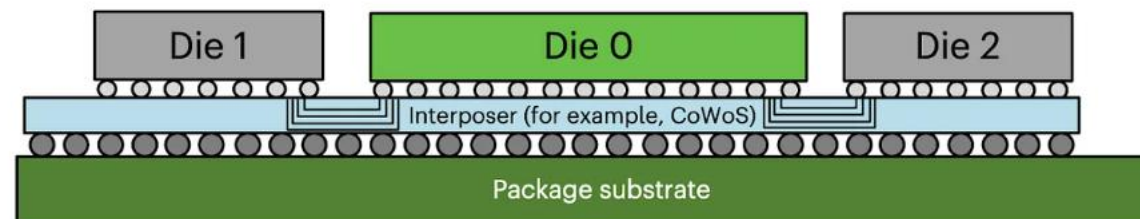
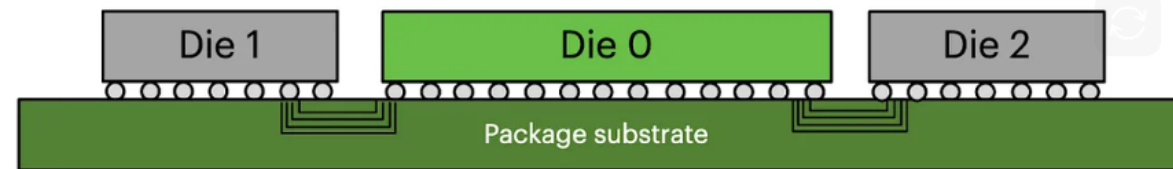
- Replicate existing designs rather than re-design larger single-cores



And also chiplets...



Standard
package





Fixed Costs

- For new chip design
 - Design & verification: ~\$100M (500 person-years @ \$200K per)
 - Amortized over “proliferations”, e.g., Core i3, i5, i7 variants
- For new (smaller) technology generation
 - ~\$3B for a new fab (future fabs probably a lot more)
 - Amortized over multiple designs (Intel’s tick/tock)
 - Amortized by “rent” from companies that don’t fab themselves
 - Xilinx invented this approach
 - Eg. Qualcomm, Apple, NVIDIA, Broadcom, MediaTek, AMD
- Moore’s Law generally increases startup cost
 - More expensive fabrication equipment
 - More complex chips take longer to design and verify

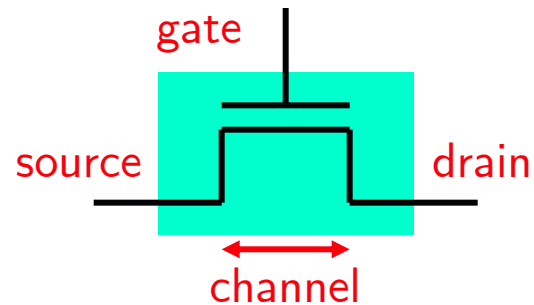


Technology Scaling



Moore's Law: Technology Scaling

- **Moore's Law**: aka “technology scaling”
 - Continued miniaturization (esp. reduction in channel length)
 - + Improves switching speed, area(cost)/transistor
 - Reduces transistor reliability
 - eg: DRAM density (transistors/area) doubles every 18 months
- Dennard Scaling: power/transistor
 - Corollary to Moore's law that has more-or-less ended





Moore's Effect #1: Transistor Count

- Linear shrink in each dimension
 - 180nm, 130nm, 90nm, 65nm, 45nm, 32nm, ...
 - Each generation is a 1.414 linear shrink
 - Shrink each dimension (2D)
 - Results in 2x more transistors (1.414×1.414)
- Reduces cost per transistor
- More transistors can increase performance
 - Job of a computer architect: use the ever-increasing number of transistors
 - Examples: caches, exploiting parallelism at all levels



Moore's Effect #2: RC Delay

- **First-order: speed scales proportional to gate length**
 - Has provided much of the performance gains in the past
- Scaling helps wire and gate delays in some ways...
 - + Transistors become shorter (Resistance↓), narrower (Capacitance↓)
 - + Wires become shorter (Length↓ → Resistance↓)
 - + Wire “surface areas” become smaller (Capacitance↓)
- Hurts in others...
 - Transistors become narrower (Resistance↑)
 - Gate insulator thickness becomes smaller (Capacitance↑)
 - Wires becomes thinner (Resistance↑)
- What to do?
 - Take the good, use wire/transistor sizing & repeaters to counter bad
 - Exploit new materials: Aluminum → Copper, metal gate, high-K

Used to get much faster, not as much any more...

Architects role: Keep Communication local!



Moore's Effect #3: Cost

- Mixed impact on unit integrated circuit cost
 - + Either lower cost for same functionality...
 - + Or same cost for more functionality
 - Difficult to achieve high yields (defects)
 - Process variation
- Increases fixed cost
 - More expensive fabrication equipment
 - Takes longer to design, verify, and test chips
 - Being on the cutting edge is expensive!
- Transistors per dollar flattens and increases after 28nm..



Moore's Effect #4: Power

- + Technology scaling reduces power/transistor...
 - Reduced sizes and surface areas reduce capacitance (C)
 - But this effect is small now (Dennard scaling is over)
- ...but increases power density and total power
 - By increasing transistors/area and total transistors
 - Faster transistors → higher frequency → more power
 - Hotter transistors leak more (thermal runaway)
- The end of voltage scaling & “dark silicon”
 - Dark silicon: can't use all of your transistors at the same time with maximum frequency
 - Two ways to deal with this:
 1. Lower frequency, do more in parallel
 - e.g. NVIDIA H100 GPU runs at 1.5GHz
 2. More specialized units that you only operate some of the time
 - e.g. Many accelerators on Mobile SoCs



Trends in Power

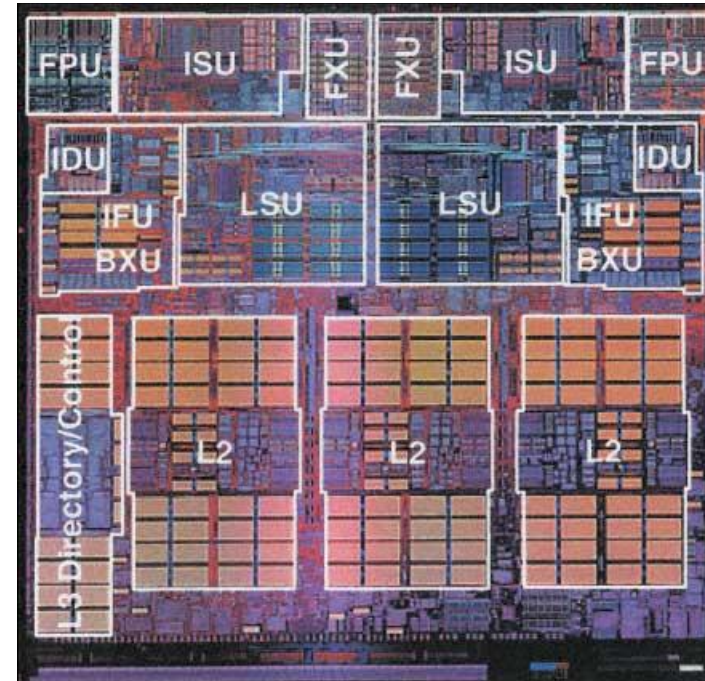
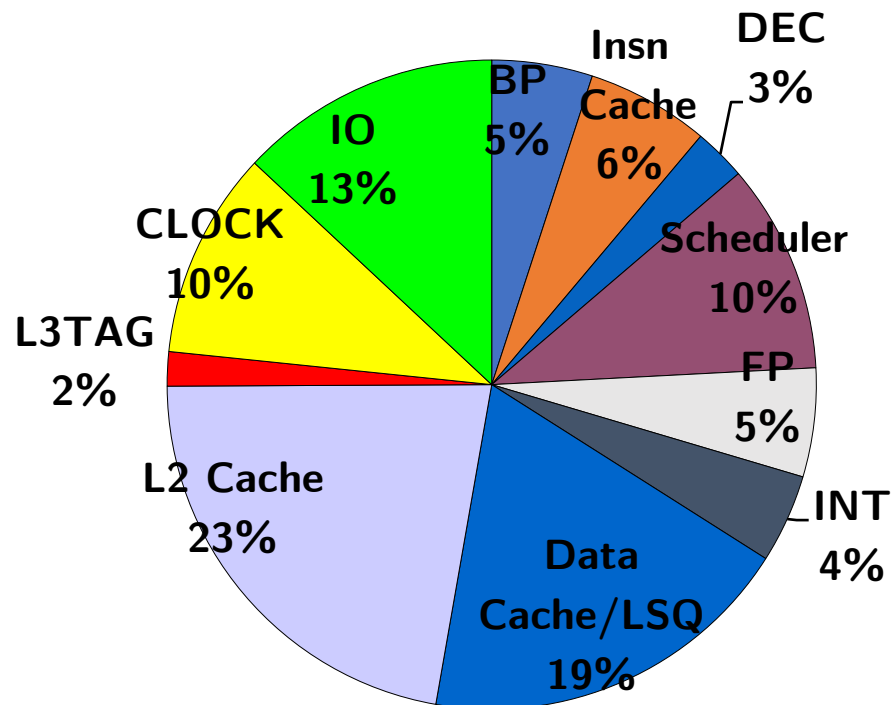
	386	Pentium	Pentium II	Pentium 4	Core2	Core i7	Cascade Lake
Year	1985	1993	1998	2001	2006	2009	2019
Technode (nm)	1500	350	180	130	65	45	14
Transistors (M)	0.3	3.1	5.5	42	291	731	8000
Voltage (V)	5	3.3	2.9	1.7	1.3	1.2	1.2
Clock (MHz)	16	66	200	1500	3000	3300	2.6-3.8K
Power (W)	1	16	35	80	75	130	100-400
Peak MIPS	6	132	600	4500	24000	52,800	582,400
MIPS/W	6	8	17	56	320	406	1,456

- Supply voltage decreasing over time
 - But “voltage scaling” is perhaps reaching its limits
- Emphasis on power starting around 2000
 - Resulting in slower frequency increases
 - Also note number of cores increasing (4 in Core i7, 28 for Casc. Lake)



Processor Power Breakdown

- Power breakdown for IBM POWER4
 - Two 4-way superscalar, 2-way multi-threaded cores, 1.5MB L2
 - Big power components are L2, data cache, scheduler, clock, I/O
 - Implications on “complicated” versus “simple” cores





Moore's Effect #5: Psychological / Economic

- **Moore's Curve:** common interpretation of Moore's Law
 - "CPU performance doubles every 18 months"
 - Self fulfilling prophecy: 2X every 18 months is $\sim 1\%$ per week
 - Q: Would you add a feature that improved performance 20% if it would delay the chip 8 months?
 - Processors under Moore's Curve (arrive too late) fail spectacularly
 - E.g., Intel's Itanium, Sun's Millennium



Moore's Law in the Future

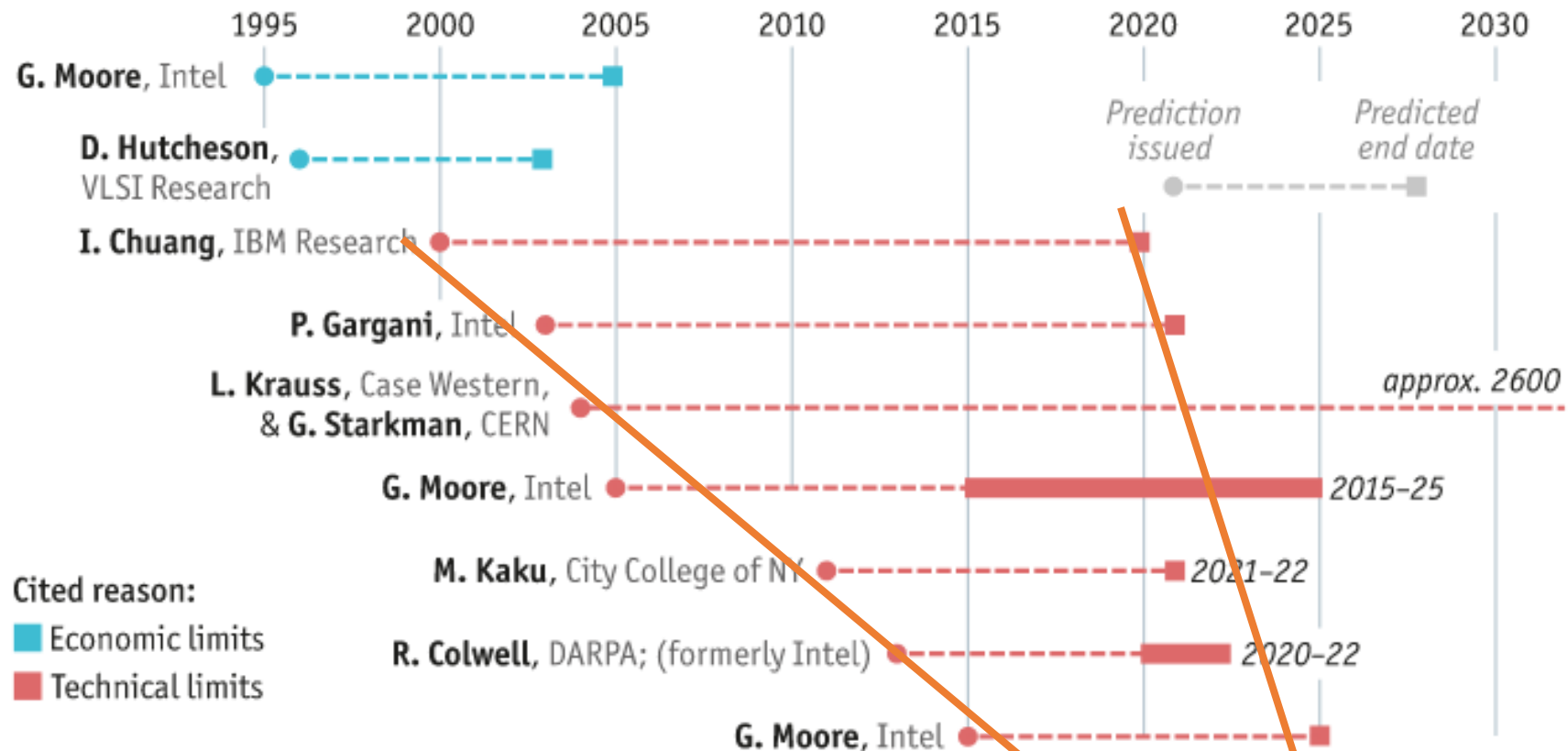
- Won't last forever, approaching physical limits
 - "If something must eventually stop, it can't go on forever"
 - But betting against it has proved foolish in the past
 - Perhaps will "slow" rather than stop abruptly
- Transistor count will likely continue to scale in short term
 - "Die stacking" is becoming main stream
 - Uses the third dimension to increase transistor count (mostly for integrating logic/memory, or stacking memory -- HBM)
 - Wafer-scale computing – don't use packages, just use wafer
 - Chiplets/Dielets – tightly integrate many small dies
- But transistor performance scaling?
 - Running into physical limits
 - Example: gate oxide is less than 10 silicon atoms thick!
 - Can't decrease it much further
 - Power is a limiting factor



Faith no Moore

Faith no Moore

Selected predictions for the end of Moore's law



Sources: Intel; press reports; *The Economist*



Technology Summary

- Has a first-order impact on computer architecture
 - Cost (die area)
 - Performance (transistor delay, wire delay)
- Most significant trends for architects (and thus CENG5410)
 - More and more transistors
 - What to do with them? → integration → **parallelism**
 - Logic is improving faster than memory & cross-chip wires
 - “Memory wall” → caches, more integration, near-data
 - Techniques to localize communication

} Rest of the term



Performance: Metrics, Means, and Mistakes



Performance

- Two considerations:
 - **Latency (execution time)**: time to finish a fixed task
 - **Throughput (bandwidth)**: number of tasks in fixed time
 - Very different: throughput can exploit parallelism, latency cannot
 - Ex: baking bread
 - Often at odds
 - Choose definition of performance that matches your goals
- Example: move people from A $\rightarrow$ B, 10 miles
 - Car: capacity = 5, speed = 60 miles/hour
 - Bus: capacity = 60, speed = 20 miles/hour
 - Latency: **car = 10 min**, bus = 30 min
 - Throughput: car = 15 PPH (count return trip), **bus = 60 PPH**



Adding/Average Performance Numbers

- You can add latencies, but not throughput
 - $\text{Latency}(P1+P2, A) = \text{Latency}(P1, A) + \text{Latency}(P2, A)$
 - $\text{Throughput}(P1+P2, A) \neq \text{Throughput}(P1, A) + \text{Throughput}(P2, A)$
 - Need total work over total time
 - $\text{Throughput}(P1+P2, A) = 2 / [(1/\text{Throughput}(P1, A)) + (1/\text{Throughput}(P2, A))]$
 - E.g. 1 mile @ 30 miles/hour + 1 mile @ 90 miles per hour
 - Average is **not** 60 miles per hour
 - 0.03333 hours at 30 miles/hour + 0.01111 hours at 90 miles/hour
 - Average is only 45 miles/hour! (2 miles/ (0.03333 + 0.01111))
- In General:
 - **Arithmetic:** $(1/N) * \sum_{P=1..N} \text{Latency}(P)$
 - For units that are proportional to time (e.g., latency)
 - **Harmonic:** $N / \sum_{P=1..N} 1/\text{Throughput}(P)$
 - For units that are inversely proportional to time (e.g., throughput)



What mean for Execution Time?

- Good: The arithmetic mean is meaningful, i.e. it's how long the average program takes to complete.
- Bad: Optimizations favor programs that take longer...
- One solution: compare speedup instead of execution time (next slide)

		Execution Time			
		P1	P2	P3	Arithmetic Mean
Machines	A	10	9	20	13.00
	B	12	8	18	12.67
	C	4	4	43	17.00



Comparing Speedup

Speedup Relative to A				
	P1	P2	P3	Arith. Mean
A/A	1.00	1.00	1.00	1.00
B/A	0.83	1.13	1.11	1.02
C/A	2.50	2.25	0.47	1.74

Speedup Relative to B				
	P1	P2	P3	Arith. Mean
A/B	1.20	0.89	0.90	1.00
B/B	1.00	1.00	1.00	1.00
C/B	3.00	2.00	0.42	1.81

Speedup Relative to C				
	P1	P2	P3	Arith. Mean
A/C	0.40	0.44	2.15	1.00
B/C	0.33	0.50	2.39	1.07
C/C	1.00	1.00	1.00	1.00

- **Geometric mean:**
$$\sqrt[N]{\prod_{P=1..N} \text{Speedup}(P)}$$
- For unitless quantities (e.g., speedup ratios)
- Only mean where the relative merit does not depend on the reference/baseline speedup



Mean Summary

In general for means (average):

- **Arithmetic:** $(1/N) * \sum_{P=1..N} \text{Latency}(P)$
 - For units that are proportional to time (e.g., latency)
- **Harmonic:** $N / \sum_{P=1..N} 1/\text{Throughput}(P)$
 - For units that are inversely proportional to time (e.g., throughput)
- **Geometric:** $\sqrt[N]{\prod_{P=1..N} \text{Speedup}(P)}$
 - For unitless quantities (e.g., speedup ratios)



Pitfalls of Partial Performance Metrics



Iron Law of Performance

$$\frac{\text{Time}}{\text{Program}} = \frac{\text{Instructions}}{\text{Program}} \times \frac{\text{Cycles}}{\text{Instruction}} \times \frac{\text{Time}}{\text{Cycle}}$$

CPI
(microarchitecture value added)



Danger: Partial Performance Metrics

- General public often equates performance with particular architecture aspects, **neglecting CPI!**
 - ISA Bitwidth (16,32,64,128 bit, yay!)
 - **Frequency**
 - # Cores
- Which processor would you buy?
 - Processor A: $\text{CPI} = 2$, clock = 5 GHz
 - Processor B: $\text{CPI} = 1$, clock = 3 GHz
 - Probably A, but B is faster (assuming same ISA/compiler)
- Classic example
 - 800 MHz PentiumIII faster than 1 GHz Pentium4!
- **Meta-point: danger of partial performance metrics!**



MIPS Metric (ignores instr. count)

- (Micro) architects often ignore dynamic instruction count
 - Typically work in one ISA/one compiler → treat it as fixed

$$\frac{\text{Time}}{\text{Instruction}} = \cancel{\frac{\text{Instructions}}{\text{Program}}} \times \frac{\text{Cycles}}{\text{Instruction}} \times \frac{\text{Time}}{\text{Cycle}}$$

- **MIPS** (millions of instructions per second, *inverse of above*)
 - **Instructions/second * 10⁻⁶**
 - **Cycles / second**: clock frequency (in MHz)
 - Example: CPI = 2, clock = 500 MHz → 0.5 * 500 MHz = 250 MIPS
- Pitfall of MIPS: may vary inversely with actual performance
 - Compiler removes insns, program gets faster, MIPS can go down (eg. instructions that it removed were easy to run in parallel...)
- Other problems:
 - Some optimizations actually add instructions
 - Work per instruction varies (e.g., multiply vs. add, FP vs. integer)
 - ISAs are not equivalent



MFLOPS/GFLOPS/TFLOPS

- **MFLOPS**: like MIPS, but counts only FP ops, because...
 - +FP ops can't be optimized away (by compiler)
 - +FP ops have longest latencies (FP Units are larger)
 - +FP ops are same across machines
- +Or in ML, equivalent thing is multiply-accumulates...
 - +Point is to pick a unit of work that is independent of implementation
- Pitfalls:
 - Many CPU programs are “integer” – light on FP
 - Loads from memory can take much longer than even an FP divide
 - Even FP instruction sets are not equivalent
- Upshot: Neither MIPS nor MFLOPS are universal



Improving CPI

$$\frac{\text{Time}}{\text{Program}} = \frac{\text{Instructions}}{\text{Program}} \times \frac{\text{Cycles}}{\text{Instruction}} \times \frac{\text{Time}}{\text{Cycle}}$$

- This course focuses on improve CPI instead of frequency
 - Until early 2000s, clock accounts for 70%+ of performance improvement
 - Achieved via deeper pipelines
 - That has been forced to change
 - Deep pipelining is **not** power efficient
 - Frequency improvements stalled around 2010...
- **We need to look at micro-arch improvements instead!**
 - **Caching, speculation, multiple issue, out-of-order issue**
 - **Vectors, multiprocessing, more ...**
 - Gotta use those extra transistors somehow, amirite?



Performance Rules of Thumb

- Make common case fast
 - **Amdahl's Law**
 - $\text{Time}_{\text{optimized}} = \text{Time}_{\text{orig}} * \text{fraction}_x / \text{Speedup}_x + \text{Time}_{\text{orig}} * (1 - \text{fraction}_x)$
 - $\text{Speedup} = \text{Time}_{\text{orig}} / \text{Time}_{\text{optimized}}$
 - $\text{Speedup} = 1 / ((1 - \text{fraction}_x) + \text{fraction}_x / \text{Speedup}_x)$
 - Point: don't optimize 5% to the detriment of the other 95%
 - $\text{Speedup}_{\text{overall}} = 1 / ((1 - 5\%) + 5\% / \infty) \sim 1.05$
- Build a balanced system
 - Don't over-engineer capabilities that cannot be utilized
 - Try to be "bound" by the most expensive resource (if not everywhere)
- Design for actual, not peak performance
 - For actual performance X, machine capability must be $> X$