

Lecture 11

Memory Hierarchy

CS213 – Intro to Computer Systems
Branden Ghen a – Winter 2026

Slides adapted from:

St-Amour, Hardavellas, Bustamente (Northwestern), Bryant, O'Hallaron (CMU), Garcia, Weaver (UC Berkeley)

Administrivia

- Labs are still ongoing
 - Bomb Lab due today
 - Attack Lab is out now
- Reminder on Attack Lab partnership survey if you want a partner!
- Midterm exam grading should be done this afternoon/evening

Changing the focus of CS213

- So far in class we've focused on *how* computers do things
 - Represent data
 - Run instructions
- Now we're focusing on *how to improve* those things
 - Secure (last lecture plus some stuff in two weeks)
 - Efficient (today and next week)
- As we'll show today, the most important thing to speed up is memory
 - It is possible and hardware does so already
 - Software can be designed to take advantage of this

Topics remaining in CS213

- Memory System
 - Memory Hierarchy
 - Caches
 - Virtual Memory
- Applications
 - Concurrency
 - Processes
- Other Systems Introductions
 - Compilers
 - Networks

Today's Goals

- Hardware day!!
- Describe what's going on inside a processor
 - Assembly-to-transistors deep-dive
- Explore the memory systems available in modern computers
 - Understand capabilities and limitations of each
- Discuss the memory hierarchy
 - How it improves performance through *caching*

Outline

- **What's in a processor? Assembly-to-Transistors**
- Memory Technologies
- Memory Hierarchy
- Caching Concept

Assembly into machine code

```
test:
4011da  48 8d 04 7e    lea    (%rsi,%rdi,2),%rax
4011de  48 8d 04 10    lea    (%rax,%rdx,1),%rax
4011e2  48 29 f7      sub    %rsi,%rdi
4011e5  48 01 f8      add    %rdi,%rax
4011e8  48 8d 84 08 13 02 00 00    lea    0x213(%rax,%rcx,1),%rax
4011ef  c3           ret
4011f0
```

- Machine code are the numerical versions of each instruction
- Number breaks down into parts
 - Operation
 - Source
 - Destination
- immediates are stored in the instruction encoding

Representing instructions as numbers

- Why represent instructions as numbers?
 1. Everything in memory is “just a number”
 - And instructions go in memory
 2. Hardware can “decode” number to figure out what to do
 - Break number apart into bits (just like floating point)
 - Some bits pick operation
 - Some bits pick register or specify immediate

Machine code ideas

- Example:
 - ADD \$0x4351FF23, %rax
 - ADD with destination %rax translates into 0x05
 - Immediate is appended on to that
 - Machine code: 0x0523FF5143
- Number of bytes for each instruction is variable
 - 1-15 bytes depending on instruction and operands
 - So the processor reads it byte-by-byte to figure out the meaning
- Translation is tedious
 - We're not going to do it by hand, although Attack Lab touches it a bit

Computer Processor (in five easy steps)

1. Reads instruction from memory
2. Decodes it into an Operation plus Configurations
 - Immediates, Registers, Memory, etc.
3. Reads from source (based on configuration)
4. Executes that operation
5. Writes to destination (based on configuration)

These steps are relatively easy (we'll skip them)

1. Reads instruction from memory
- 2.
3. Reads from source (based on configuration)
- 4.
5. Writes to destination (based on configuration)

This is extremely complicated for x86-64 (skip it too)

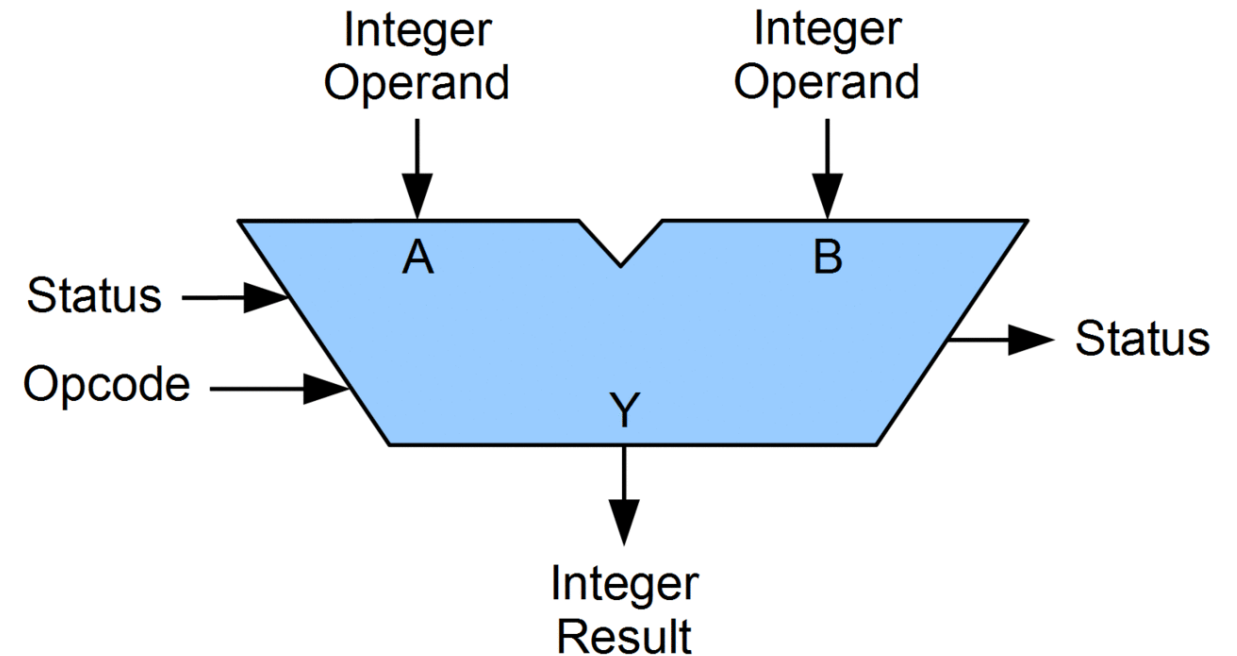
2. Decodes it into an Operation plus Configurations
 - Immediates, Registers, Memory, etc.

We can talk about what execution means though!

4. Executes that operation

Arithmetic Logic Unit (ALU)

- Piece of hardware
- Takes in two operands
 - Source and Destination *values*
- Takes in an Opcode
 - Which operation to run
- Performs operation and outputs result

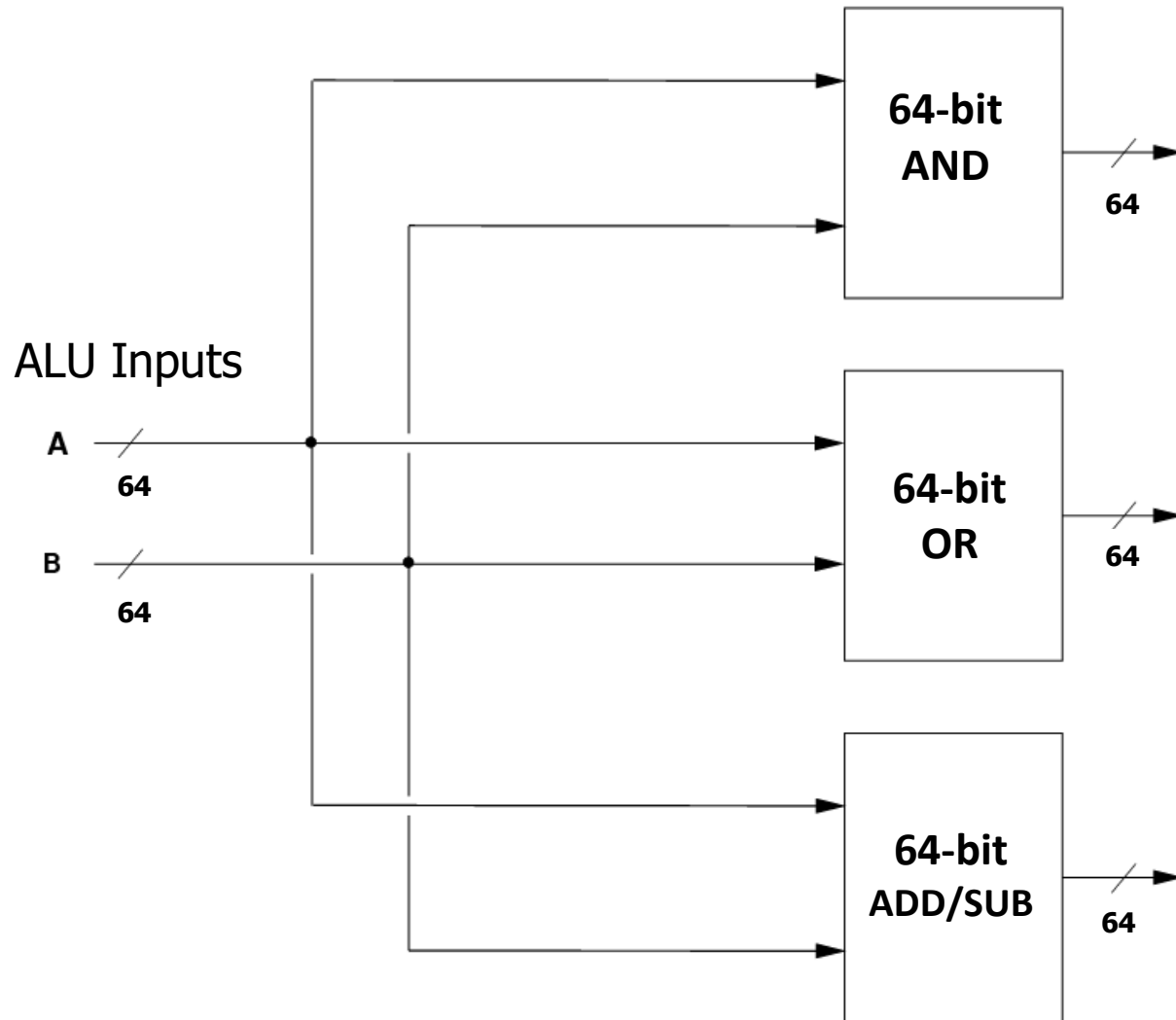


What can an ALU do?

- All the basic arithmetic operations
 - Add
 - Subtract
 - Bitwise And
 - Bitwise Or
 - Bitwise Xor
 - Arithmetic Shift Right
 - Logical Shift Right
 - Logical Shift Left
- Complex operations are separate hardware
 - Multiply, Divide, Anything floating point

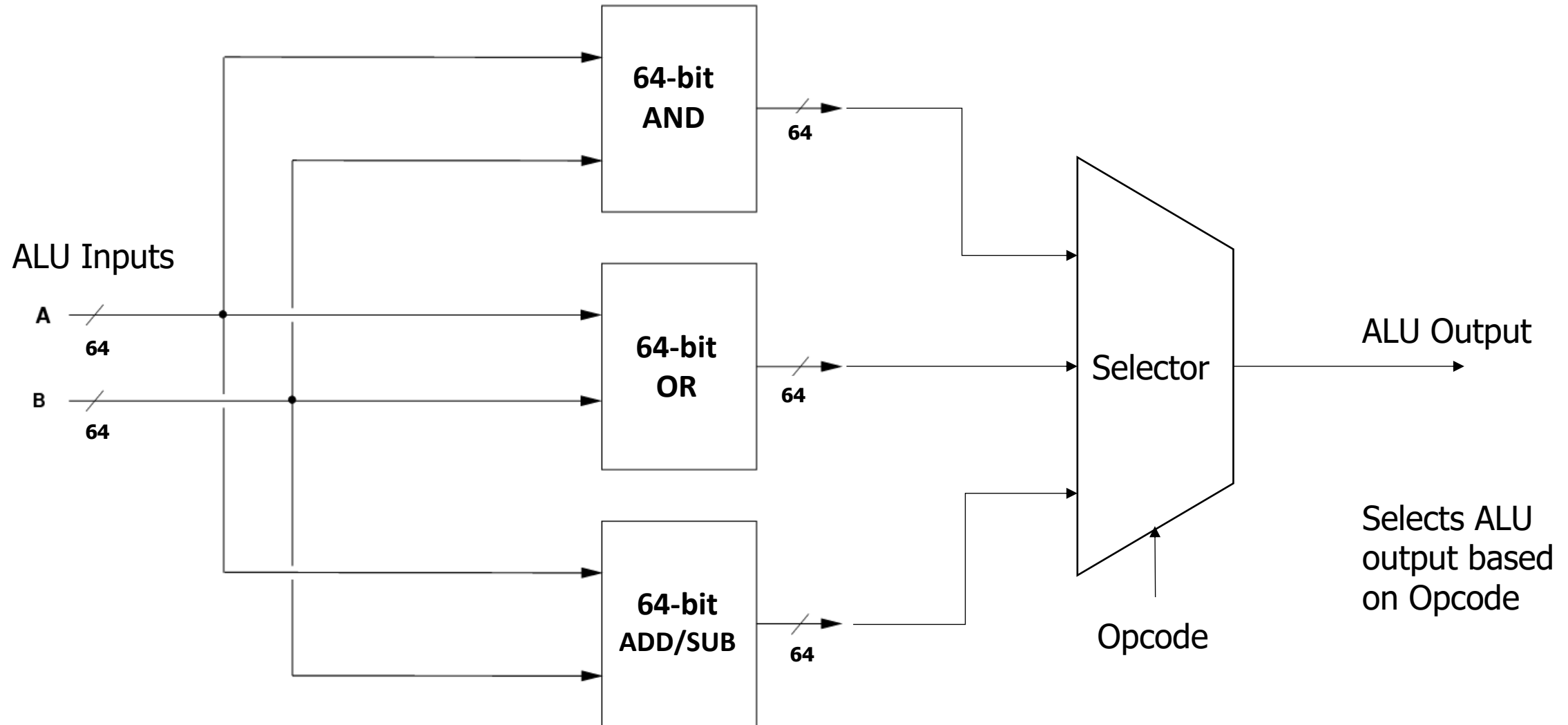
Let's zoom in

Inside an ALU



- Input values go into separate hardware blocks for each operation
- Every operation occurs in parallel, simultaneously
 - We are in hardware so this doesn't take any additional time

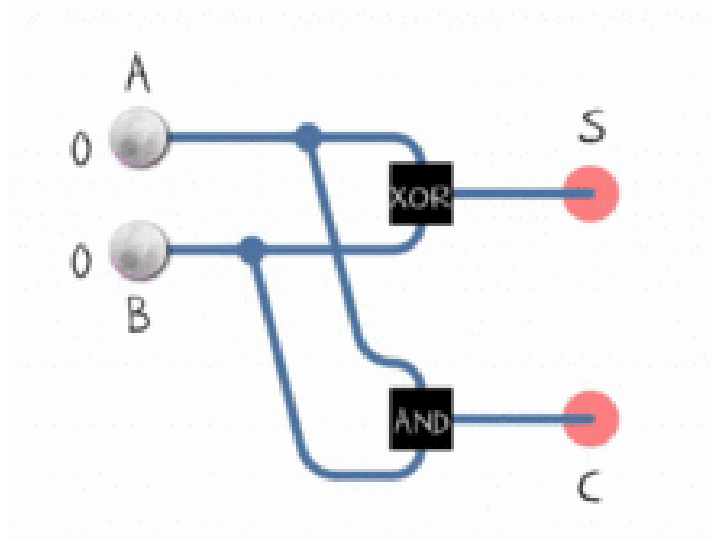
Inside an ALU – selecting the correct output

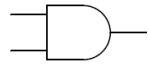
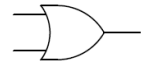
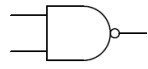
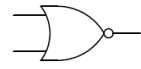
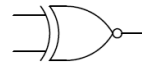


Let's zoom in

How is an ALU made?

- All of those arithmetic operations can be broken down into a series of 1-bit Boolean operations
 - Add is XOR for result + AND for carry
 - Subtract is Flip bits (NOT), Add one (XOR + AND), then Add (XOR + AND)

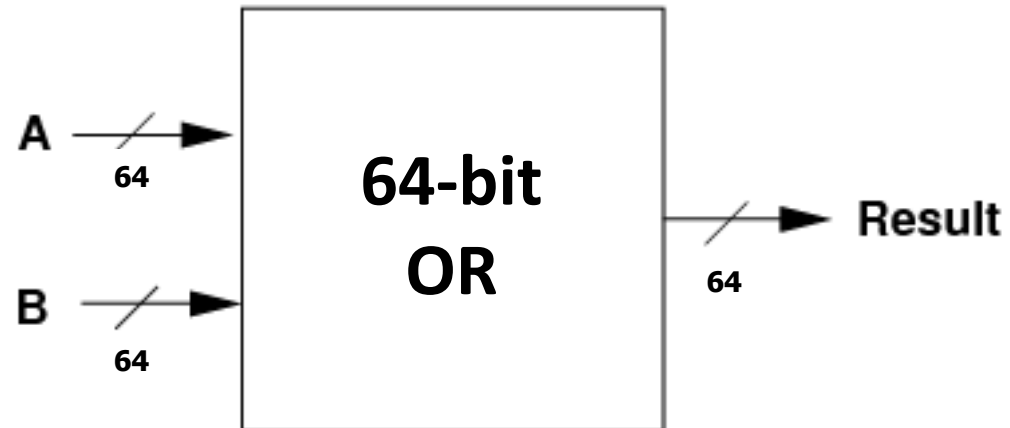
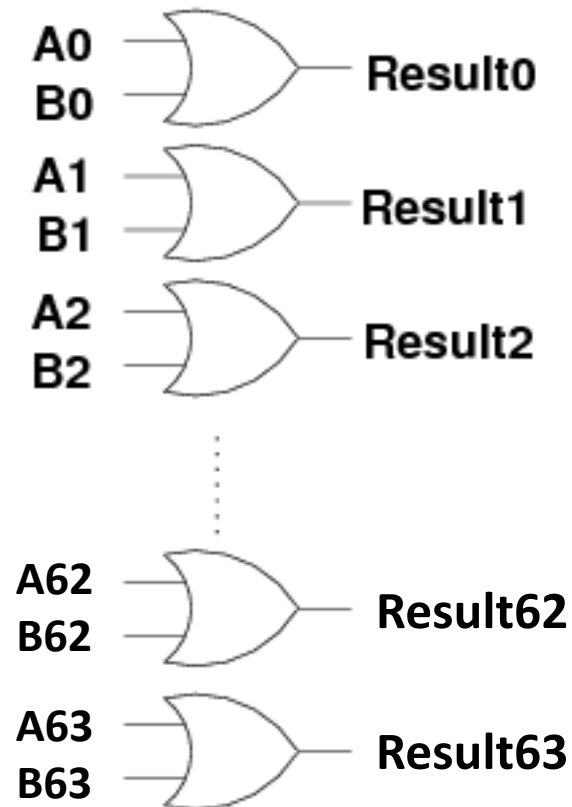


																																															
AND	OR	XOR																																													
<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>0</td></tr><tr><td>0</td><td>1</td><td>0</td></tr><tr><td>1</td><td>0</td><td>0</td></tr><tr><td>1</td><td>1</td><td>1</td></tr></table>	A	B	Output	0	0	0	0	1	0	1	0	0	1	1	1	<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>0</td></tr><tr><td>0</td><td>1</td><td>1</td></tr><tr><td>1</td><td>0</td><td>1</td></tr><tr><td>1</td><td>1</td><td>1</td></tr></table>	A	B	Output	0	0	0	0	1	1	1	0	1	1	1	1	<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>0</td></tr><tr><td>0</td><td>1</td><td>1</td></tr><tr><td>1</td><td>0</td><td>1</td></tr><tr><td>1</td><td>1</td><td>0</td></tr></table>	A	B	Output	0	0	0	0	1	1	1	0	1	1	1	0
A	B	Output																																													
0	0	0																																													
0	1	0																																													
1	0	0																																													
1	1	1																																													
A	B	Output																																													
0	0	0																																													
0	1	1																																													
1	0	1																																													
1	1	1																																													
A	B	Output																																													
0	0	0																																													
0	1	1																																													
1	0	1																																													
1	1	0																																													
																																															
NAND	NOR	XNOR																																													
<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>1</td></tr><tr><td>0</td><td>1</td><td>1</td></tr><tr><td>1</td><td>0</td><td>1</td></tr><tr><td>1</td><td>1</td><td>0</td></tr></table>	A	B	Output	0	0	1	0	1	1	1	0	1	1	1	0	<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>1</td></tr><tr><td>0</td><td>1</td><td>0</td></tr><tr><td>1</td><td>0</td><td>0</td></tr><tr><td>1</td><td>1</td><td>0</td></tr></table>	A	B	Output	0	0	1	0	1	0	1	0	0	1	1	0	<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>1</td></tr><tr><td>0</td><td>1</td><td>0</td></tr><tr><td>1</td><td>0</td><td>0</td></tr><tr><td>1</td><td>1</td><td>1</td></tr></table>	A	B	Output	0	0	1	0	1	0	1	0	0	1	1	1
A	B	Output																																													
0	0	1																																													
0	1	1																																													
1	0	1																																													
1	1	0																																													
A	B	Output																																													
0	0	1																																													
0	1	0																																													
1	0	0																																													
1	1	0																																													
A	B	Output																																													
0	0	1																																													
0	1	0																																													
1	0	0																																													
1	1	1																																													

- And/Or/Xor are just their respective operations
- Shifts are just move the bits around (simple in hardware, just move wires)

64-bit OR operation

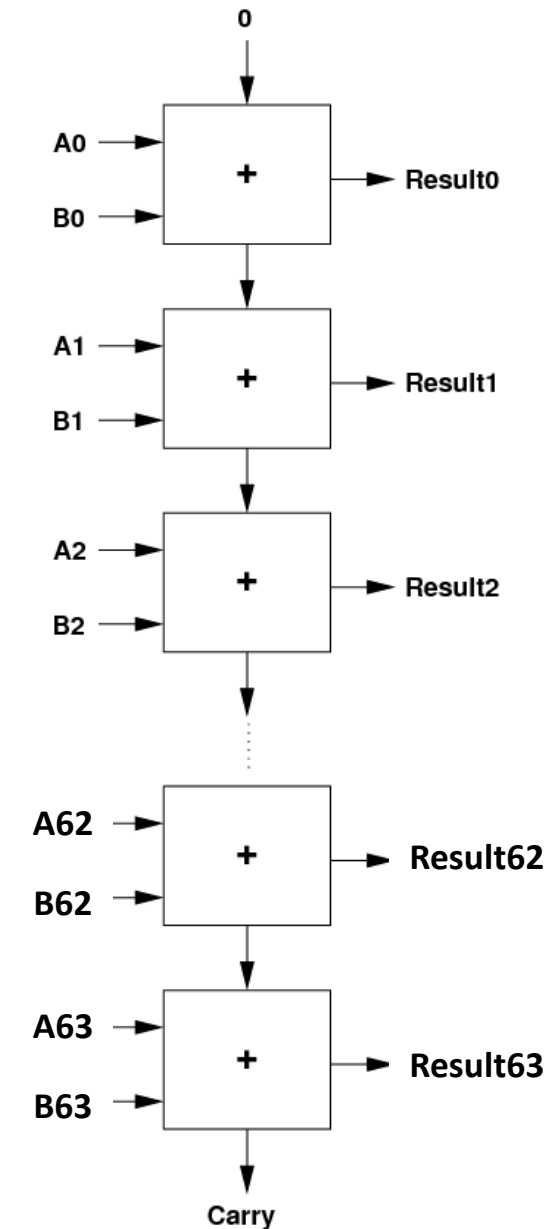
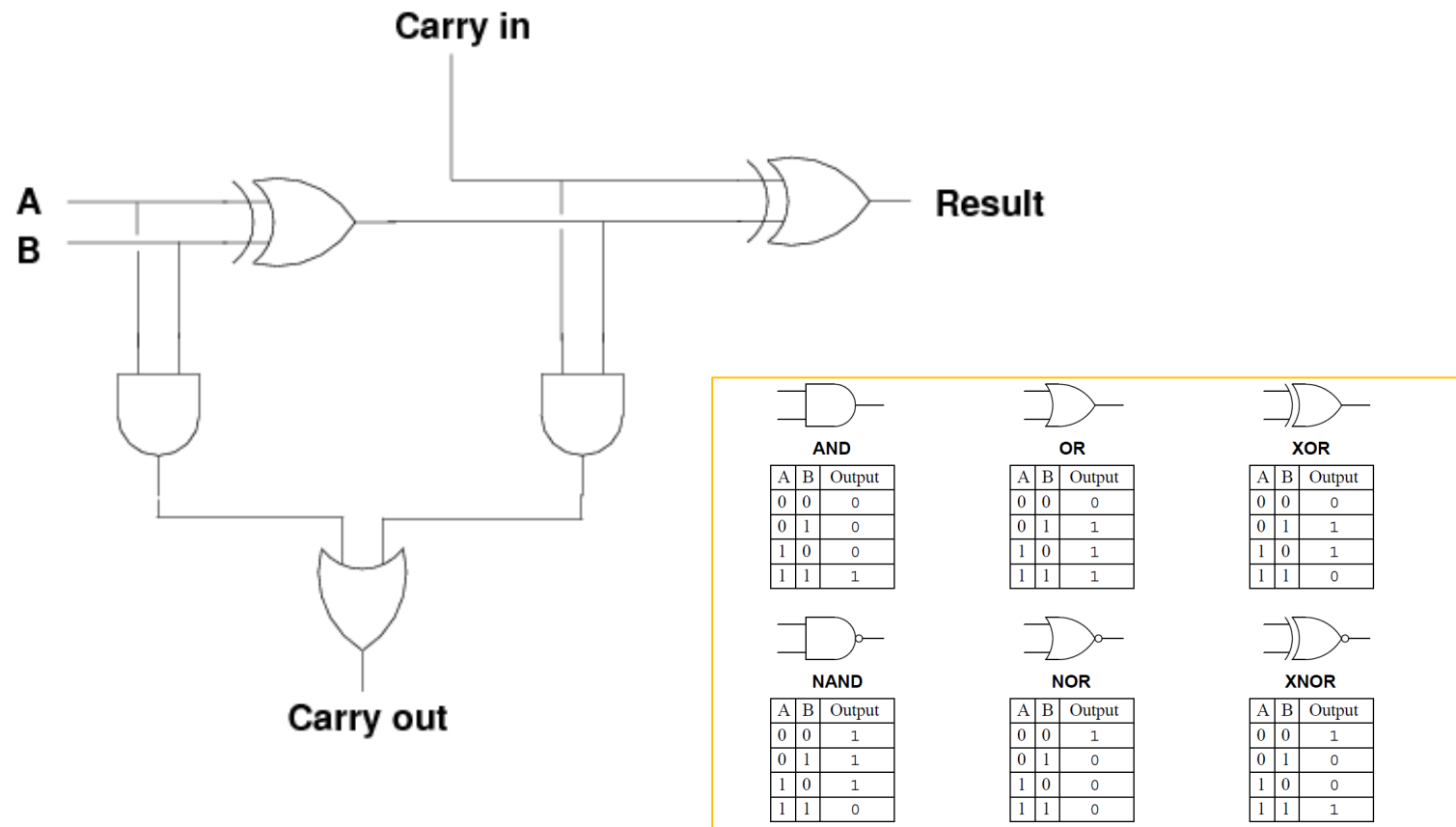
- Perform OR operation on each individual bit
 - Pictured is a series of 1-bit OR gates



																																															
AND	OR	XOR																																													
<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>0</td></tr><tr><td>0</td><td>1</td><td>0</td></tr><tr><td>1</td><td>0</td><td>0</td></tr><tr><td>1</td><td>1</td><td>1</td></tr></table>	A	B	Output	0	0	0	0	1	0	1	0	0	1	1	1	<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>0</td></tr><tr><td>0</td><td>1</td><td>1</td></tr><tr><td>1</td><td>0</td><td>1</td></tr><tr><td>1</td><td>1</td><td>1</td></tr></table>	A	B	Output	0	0	0	0	1	1	1	0	1	1	1	1	<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>0</td></tr><tr><td>0</td><td>1</td><td>1</td></tr><tr><td>1</td><td>0</td><td>1</td></tr><tr><td>1</td><td>1</td><td>0</td></tr></table>	A	B	Output	0	0	0	0	1	1	1	0	1	1	1	0
A	B	Output																																													
0	0	0																																													
0	1	0																																													
1	0	0																																													
1	1	1																																													
A	B	Output																																													
0	0	0																																													
0	1	1																																													
1	0	1																																													
1	1	1																																													
A	B	Output																																													
0	0	0																																													
0	1	1																																													
1	0	1																																													
1	1	0																																													
																																															
NAND	NOR	XNOR																																													
<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>1</td></tr><tr><td>0</td><td>1</td><td>1</td></tr><tr><td>1</td><td>0</td><td>1</td></tr><tr><td>1</td><td>1</td><td>0</td></tr></table>	A	B	Output	0	0	1	0	1	1	1	0	1	1	1	0	<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>1</td></tr><tr><td>0</td><td>1</td><td>0</td></tr><tr><td>1</td><td>0</td><td>0</td></tr><tr><td>1</td><td>1</td><td>0</td></tr></table>	A	B	Output	0	0	1	0	1	0	1	0	0	1	1	0	<table><tr><th>A</th><th>B</th><th>Output</th></tr><tr><td>0</td><td>0</td><td>1</td></tr><tr><td>0</td><td>1</td><td>0</td></tr><tr><td>1</td><td>0</td><td>0</td></tr><tr><td>1</td><td>1</td><td>1</td></tr></table>	A	B	Output	0	0	1	0	1	0	1	0	0	1	1	1
A	B	Output																																													
0	0	1																																													
0	1	1																																													
1	0	1																																													
1	1	0																																													
A	B	Output																																													
0	0	1																																													
0	1	0																																													
1	0	0																																													
1	1	0																																													
A	B	Output																																													
0	0	1																																													
0	1	0																																													
1	0	0																																													
1	1	1																																													

64-bit ADD operation

- Below is the 1-bit version with carry-in/out
 - Two 1-bit AND, two 1-bit XOR, one 1-bit OR
 - Repeat 64 times, connecting carries together

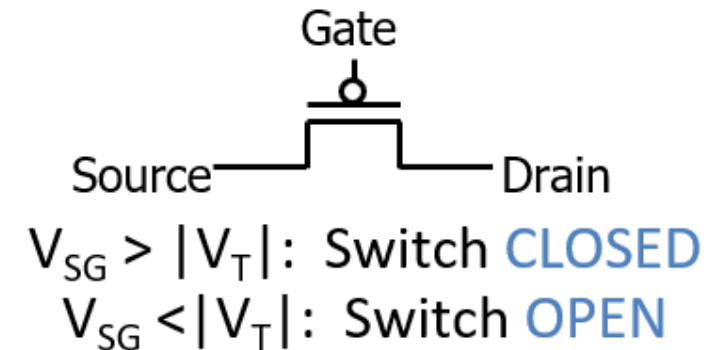
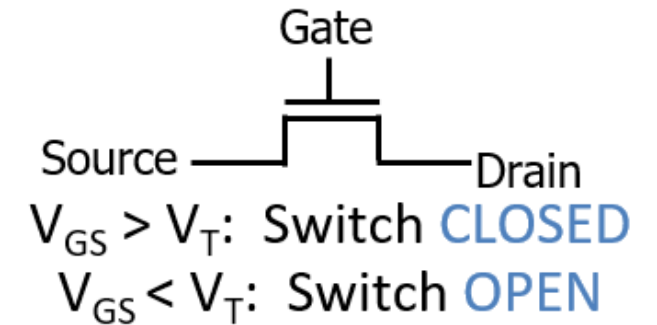
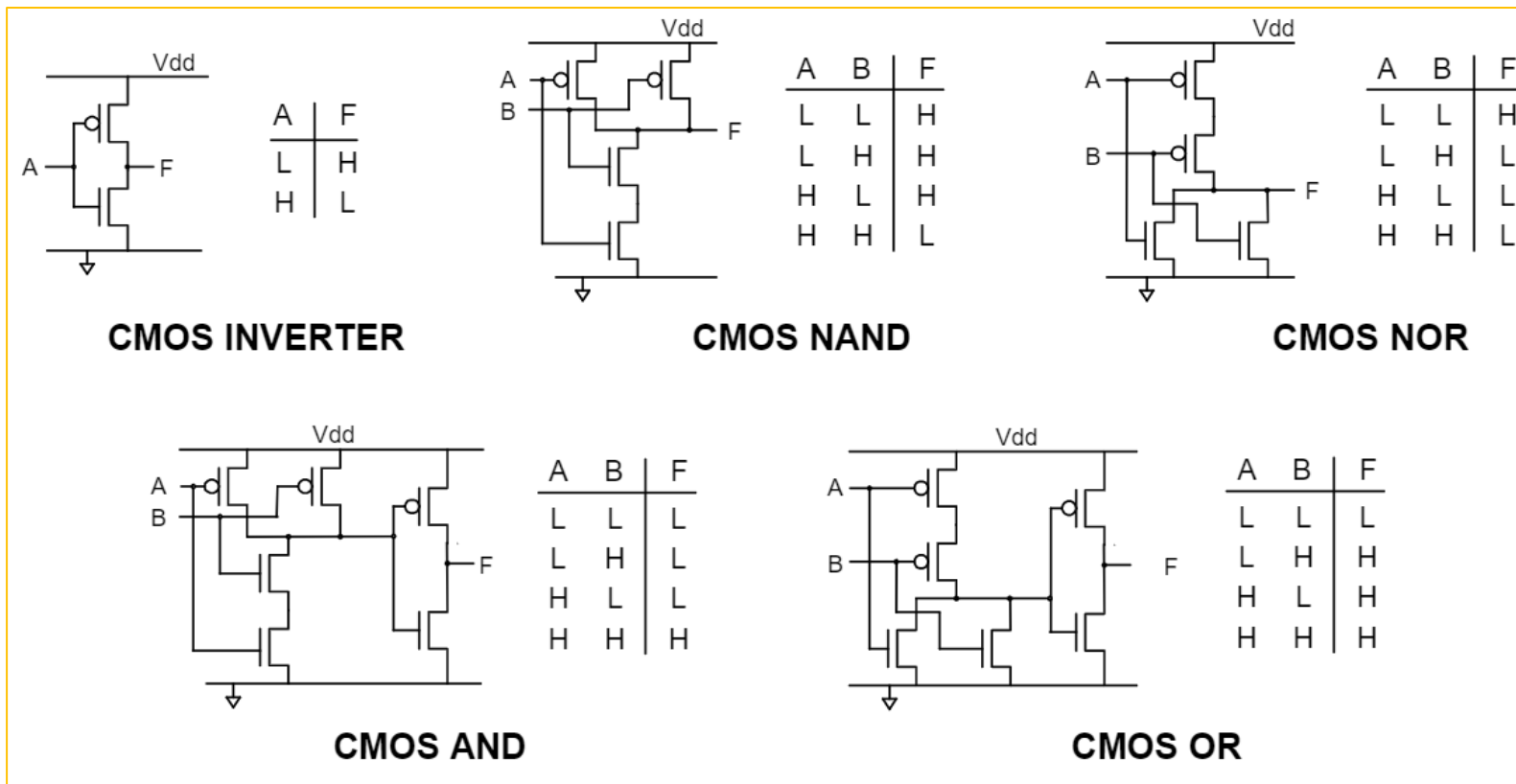


Let's zoom in

Logic gates can be created with transistors

- CMOS implementation of logic gates
 - Complementary Metal-Oxide Semiconductor

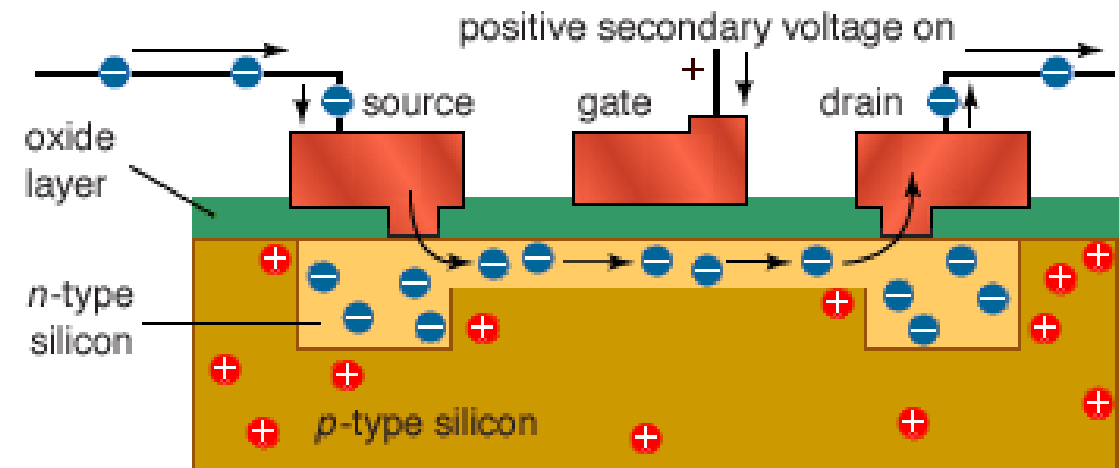
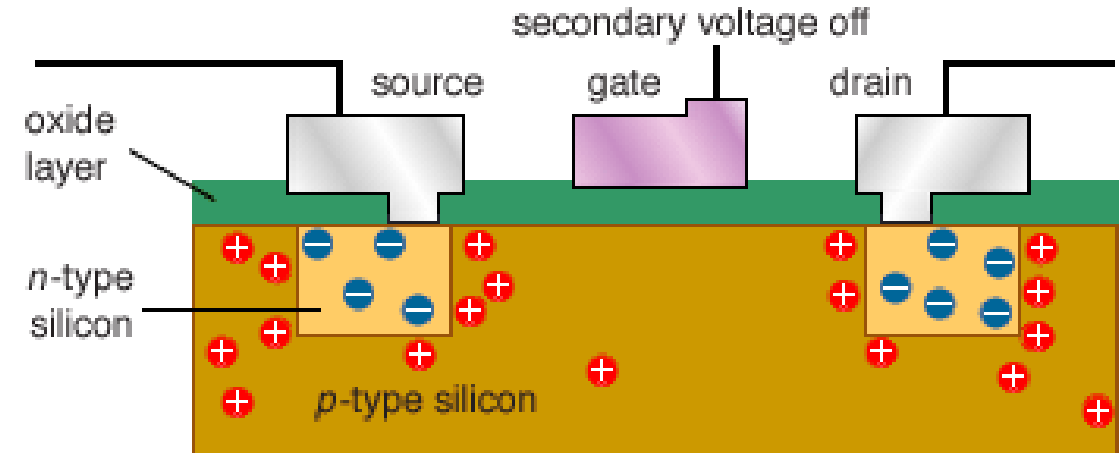
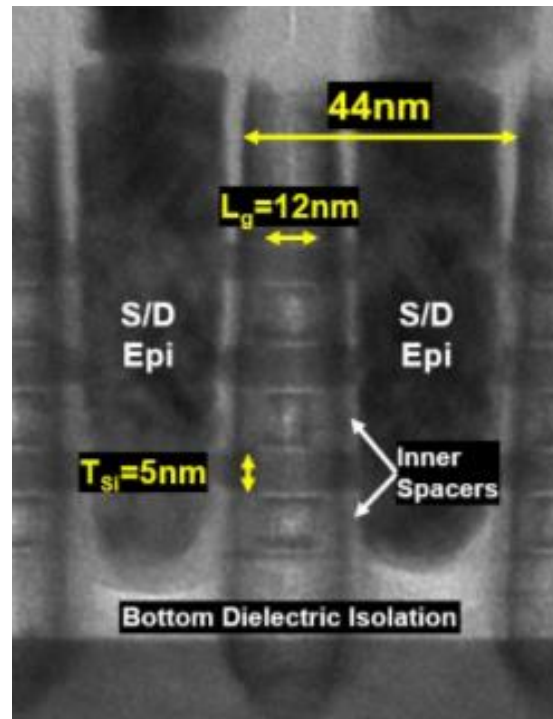
Transistors are just on/off switches



Let's zoom in

Transistors are made out of silicon and other materials

- Turning gate on/off causes source and drain to connect or disconnect
 - Acts as a switch
- We can make very small transistors
 - Handfuls of nanometers in size



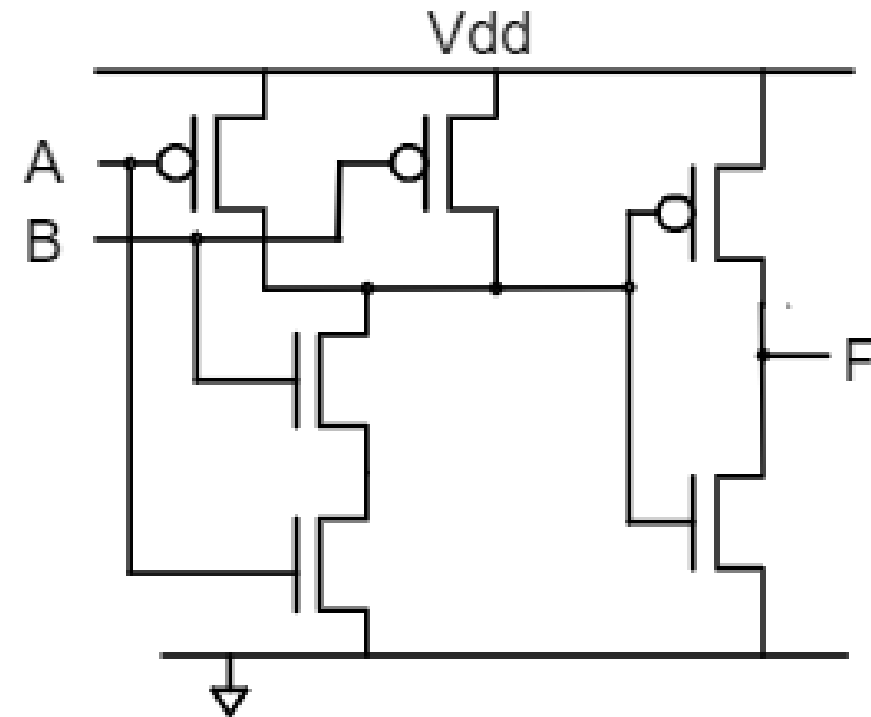
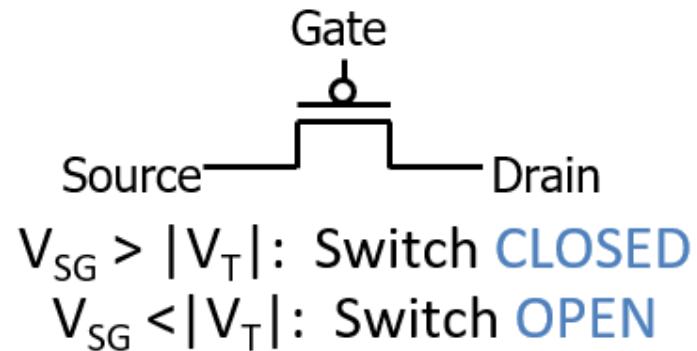
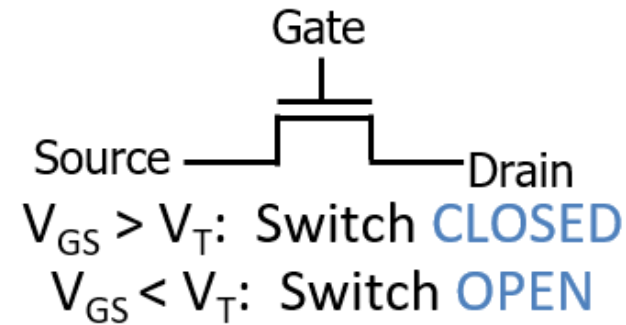
© 2004 Encyclopædia Britannica, Inc.

That's the bottom

Zooming out again

~6 transistors

- Transistors make logic gates

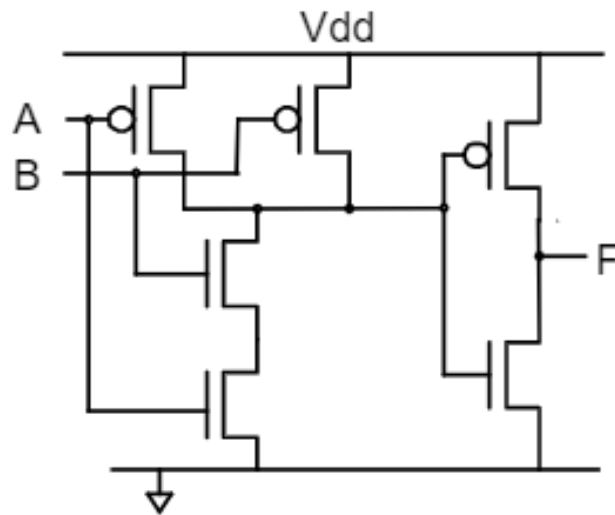


1-bit AND gate

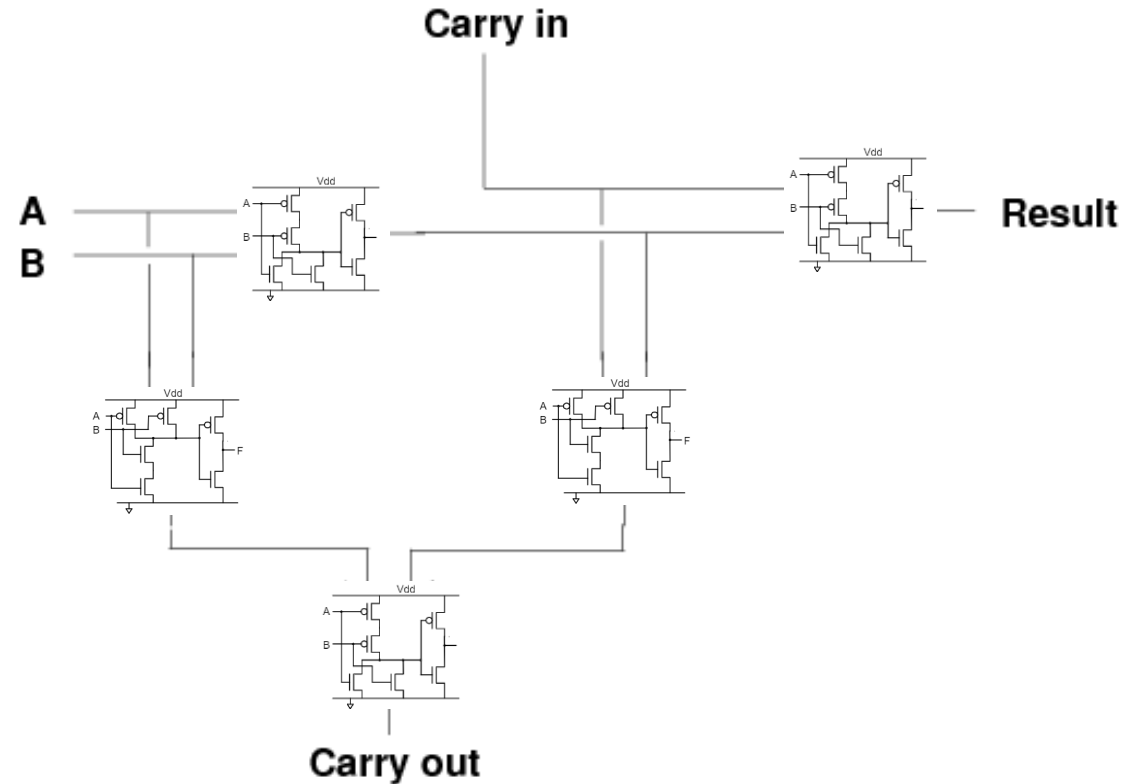
Zooming out again

~30 transistors

- Logic gates make operations



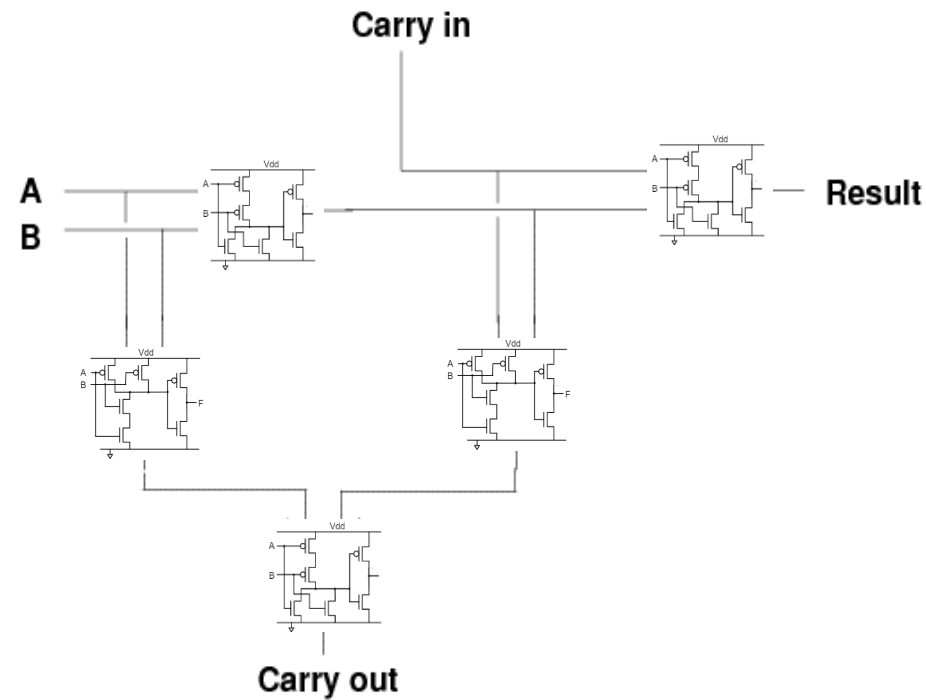
1-bit AND gate



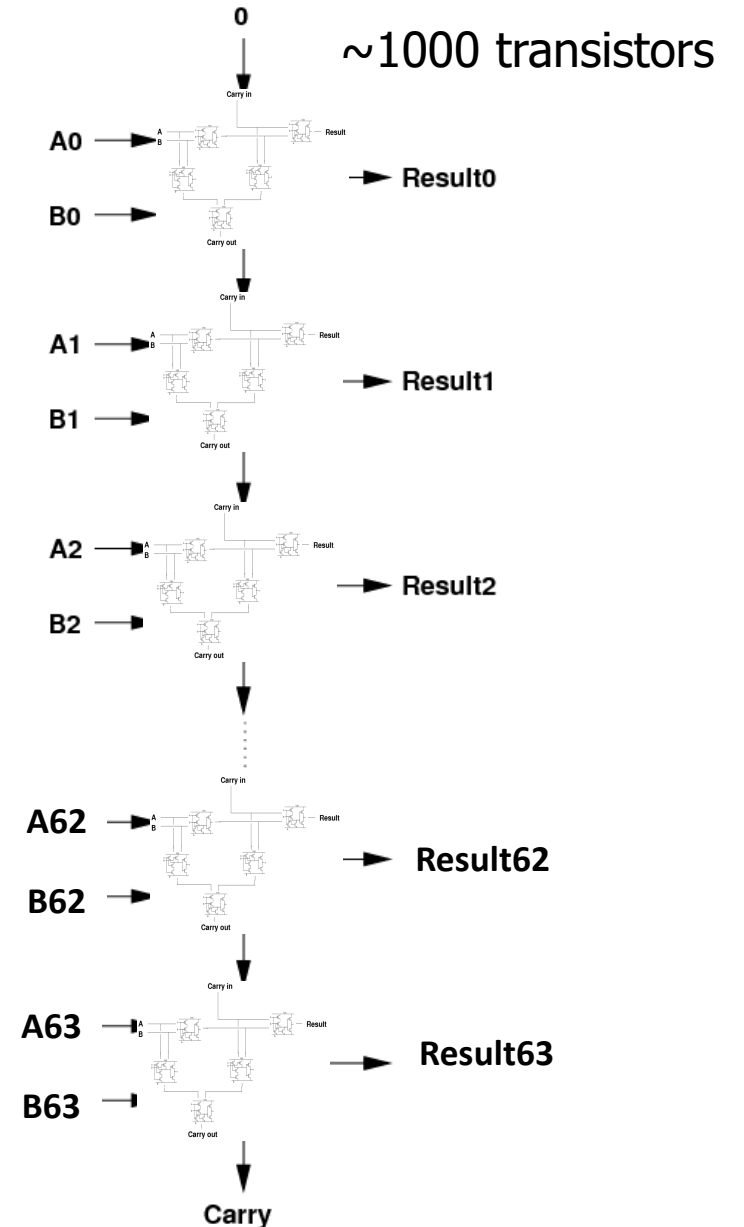
1-bit ADD operation

Zooming out again

- 1-bit operations make 64-bit operations



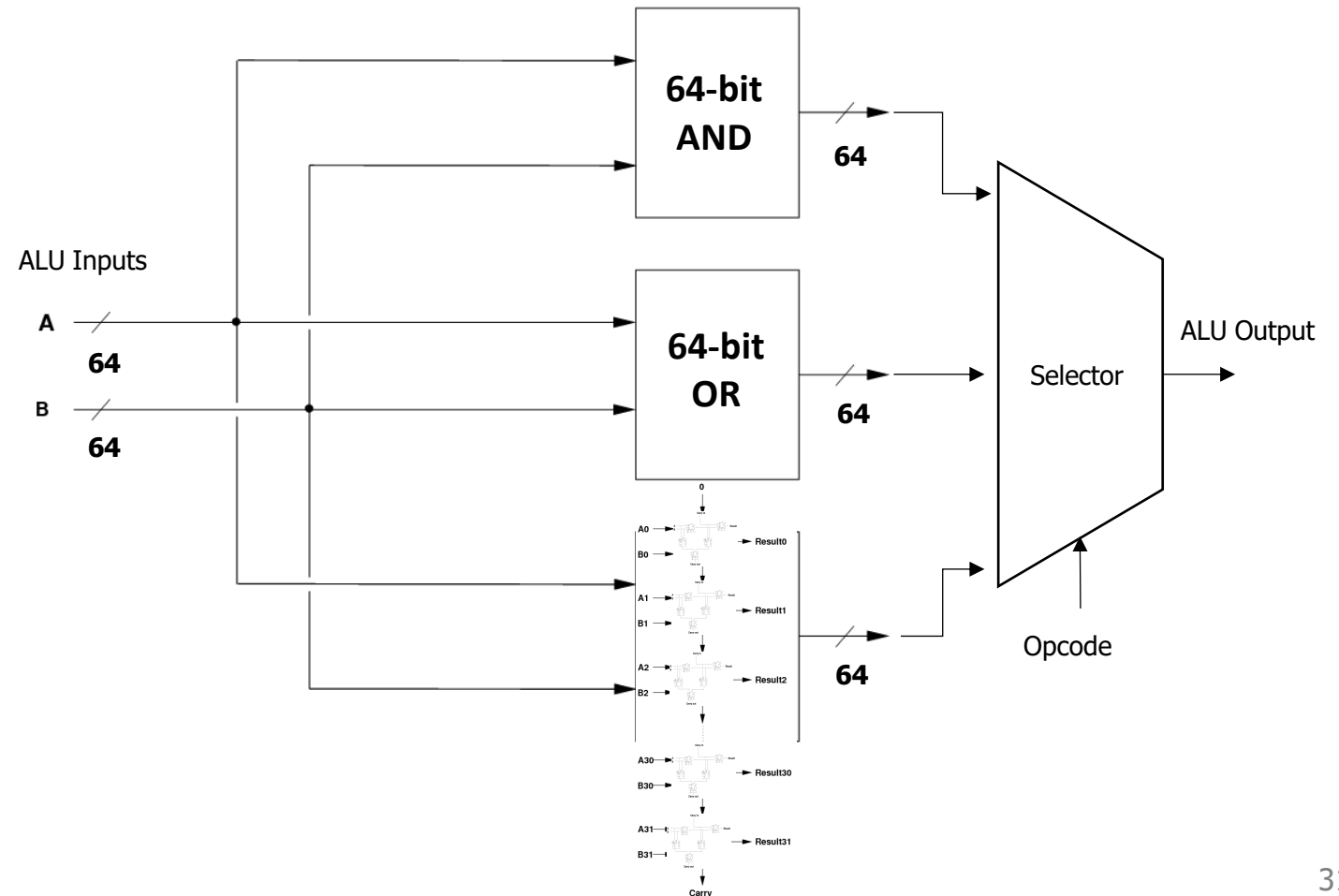
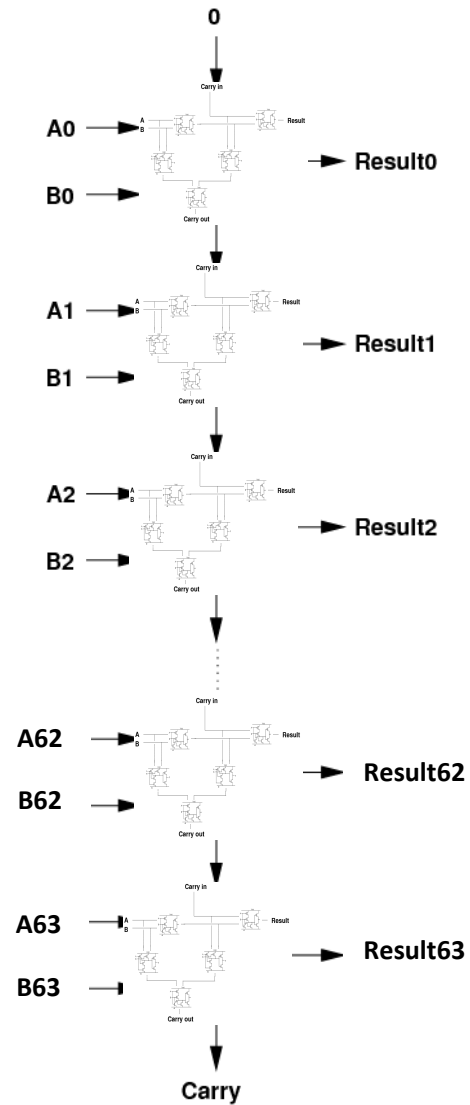
1-bit ADD operation



Zooming out again

~20K transistors

- Operations make an ALU



ALU allows us to execute operations

1. Reads instruction from memory
2. Decodes it into an Operation plus Configurations
 - Immediates, Registers, Memory, etc.
3. Reads from source (based on configuration)
- 4. Executes that operation**
5. Writes to destination (based on configuration)

All the way back to software

```
test:
    48 8d 04 7e
4011da    lea    (%rsi,%rdi,2),%rax
        48 8d 04 10
4011de    lea    (%rax,%rdx,1),%rax
        48 29 f7
4011e2    sub    %rsi,%rdi
        48 01 f8
4011e5    add    %rdi,%rax
        48 8d 84 08 13 02 00 00
4011e8    lea    0x213(%rax,%rcx,1),%rax
        c3
4011f0    ret
```

- Machine code specifies what should be executed
- Assembly translates into machine code
- C compiles into assembly

Break + Open Question

- We hit roughly 20 thousand transistors for an ALU, but a modern processor has more than 1 billion transistors.

Where are all the rest going?

Break + Open Question

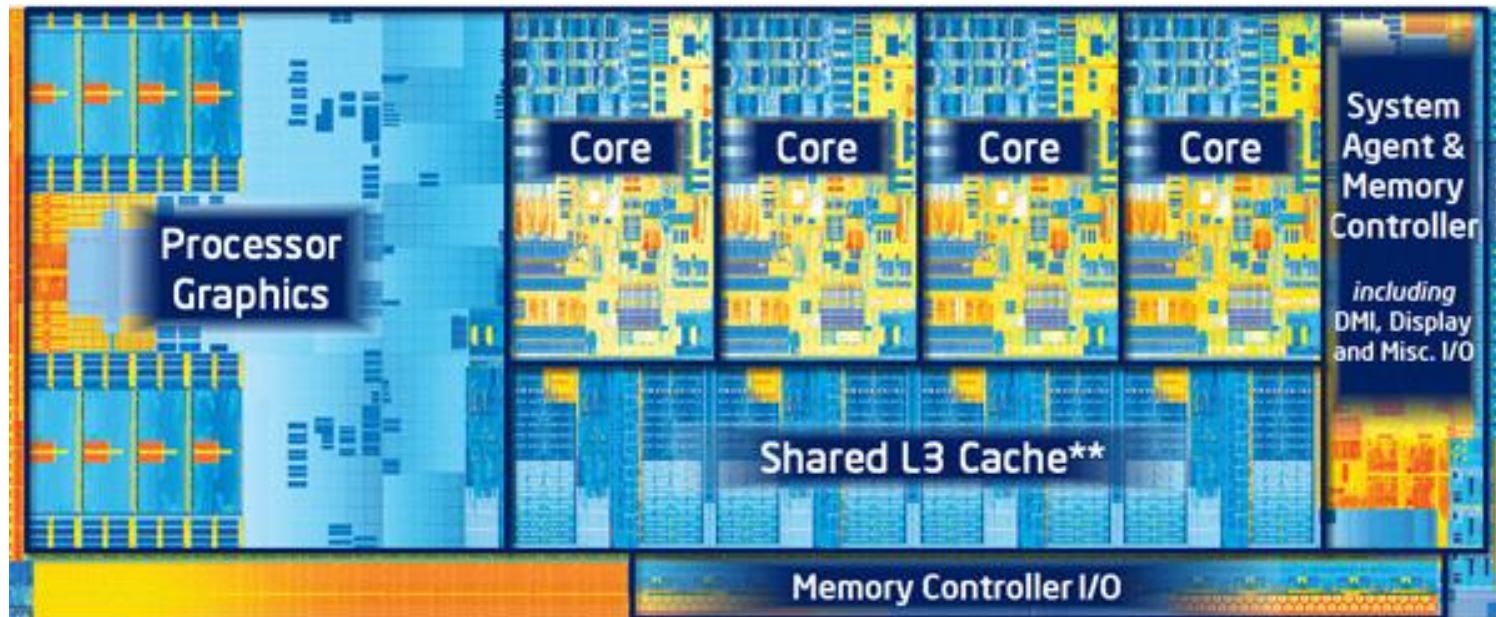
- We hit roughly 20 thousand transistors for an ALU, but a modern processor has more than 1 billion transistors.

Where are all the rest going?

- Decoding instructions
- Multiple processor cores (each with multiple ALUs and decoders)
- Graphics
- Other peripheral devices
 - USB, HDMI, Ethernet, etc.
- Memory!
 - Registers and Caches (especially caches)

A processor is just a lot of transistors connected very carefully

- ALU plus other operations make up a Core
 - Decode logic, floating point, registers, etc.
- Multiple cores plus caches make up a Processor
 - And other stuff these days like graphics



Outline

- What's in a processor? Assembly-to-Transistors
- **Memory Technologies**
- Memory Hierarchy
- Caching Concept

Reminder: metric unit prefixes

- Going to talk about nanometers and nano/micro/milliseconds today
- Also we do frequently mention bytes, kilobytes, megabytes, gigabytes, etc.
- For perspective
 - 1 Megasecond $\approx$ 11 days
 - 1 Gigasecond $\approx$ 31 years

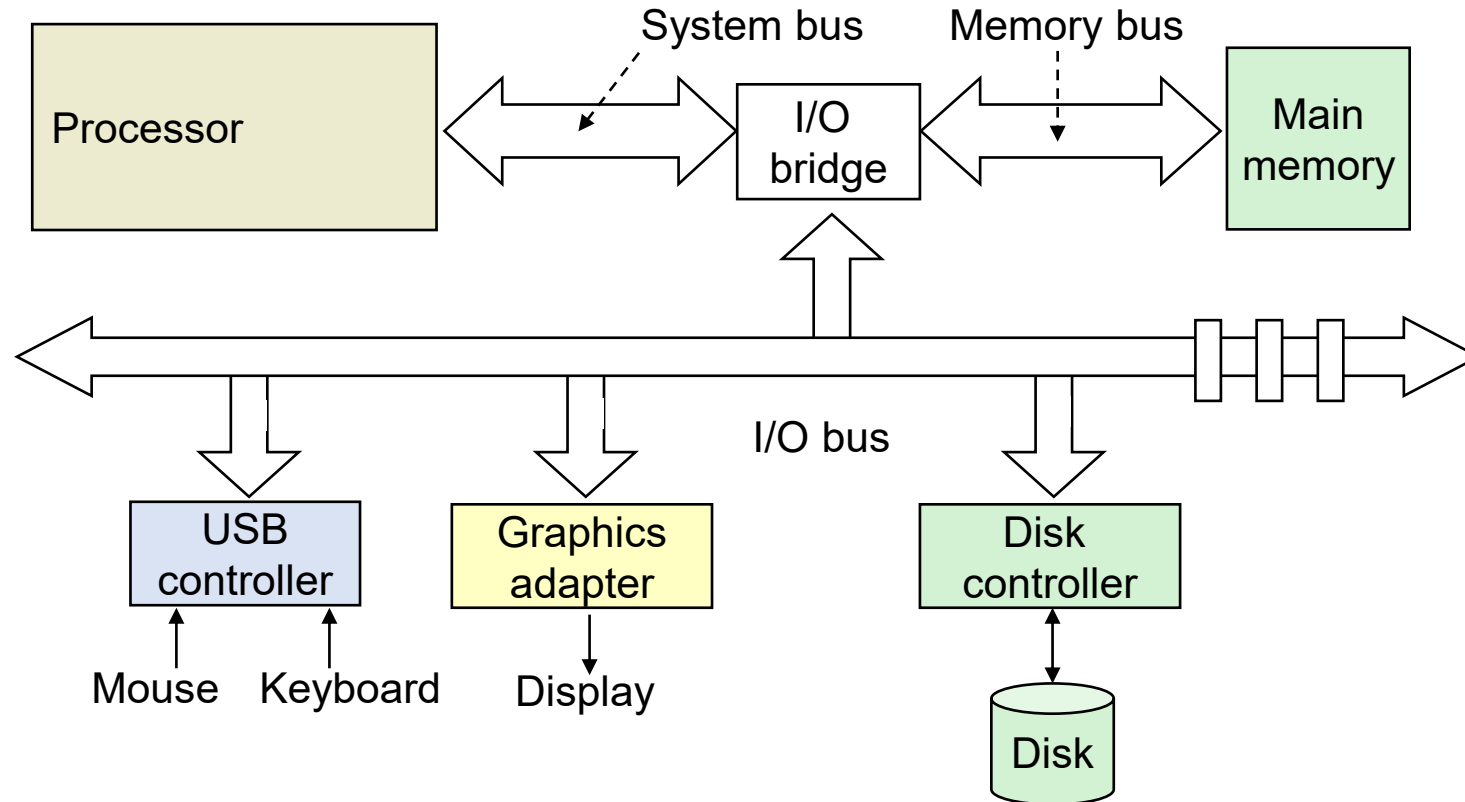
Prefix	Symbol	Power	Decimal Value
Peta	P	10^{15}	1 000 000 000 000 000
Tera	T	10^{12}	1 000 000 000 000
Giga	G	10^9	1 000 000 000
Mega	M	10^6	1 000 000
Kilo	k	10^3	1 000
-		10^0	1
milli	m	10^{-3}	0.001
micro	μ	10^{-6}	0.000 001
nano	n	10^{-9}	0.000 000 001
pico	p	10^{-12}	0.000 000 000 001

This table is something you should memorize (at least to $10^{\pm 9}$)

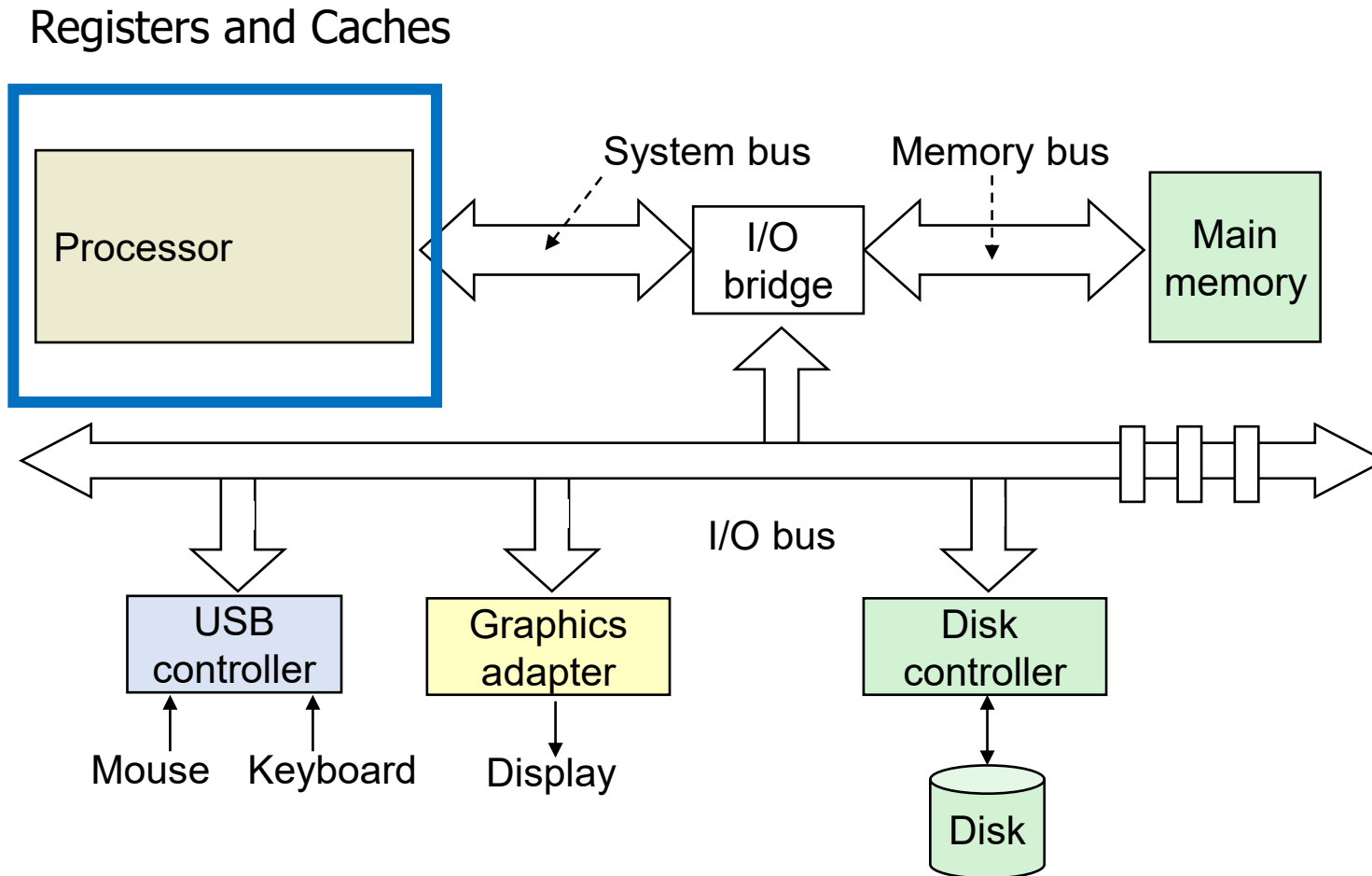
Storage in a Computer System

- Data can be stored in different places
 - Registers, caches, memory, disk, etc.
- Hugely different characteristics
 - Storage size
 - x86-64: 16 registers. 128 B total (not including FP registers, etc.)
 - Memory is measured in GB, disks in TB
 - Latency (i.e., time to access data)
 - Registers are really fast, memory less so, disk is incredibly slow
 - Cost per byte
 - Registers are really expensive, disks are really cheap, memory somewhere in the middle
- Each serves a different purpose in a system
 - We design systems to get the best of all worlds (in most cases)

Tour of computer memory



Tour of computer memory



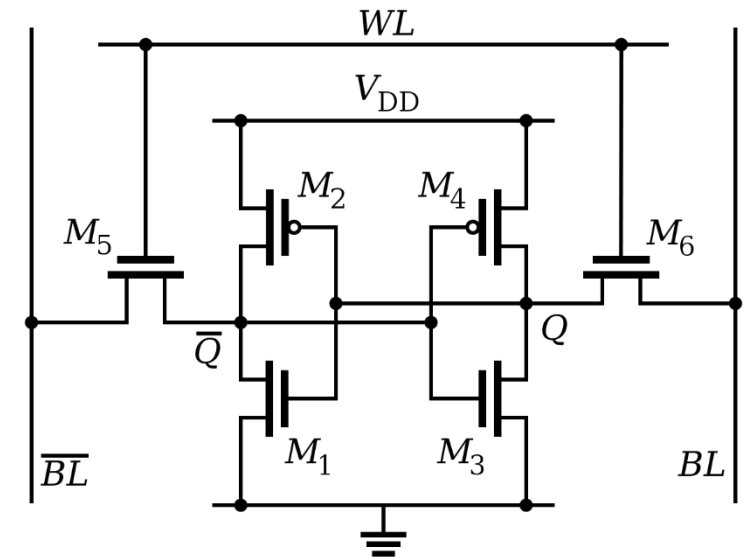
Register storage

- x86-64
 - 16 registers (general-purpose) with 8 bytes each
 - 32 registers (special-purpose) with up to 64 bytes each
 - Plus some other odds and ends (`%rip`, flags, segments, etc.)
- 128 bytes for general purpose registers
- Order 2-3 kB for everything (including floating point stuff)
- Accesses are very fast. Within a single processor cycle
 - Usually, 0.25-4 cycles per instruction

Register technology: SRAM

- Static RAM (SRAM)

- Each cell stores a bit in a bi-stable circuit, typically a six-transistor circuit
- Static – no need for periodic refreshing; keeps data while powered
- Relatively insensitive to disturbances such as electrical noise
 - Energetic particles (alpha particles, cosmic rays) can flip stored bits



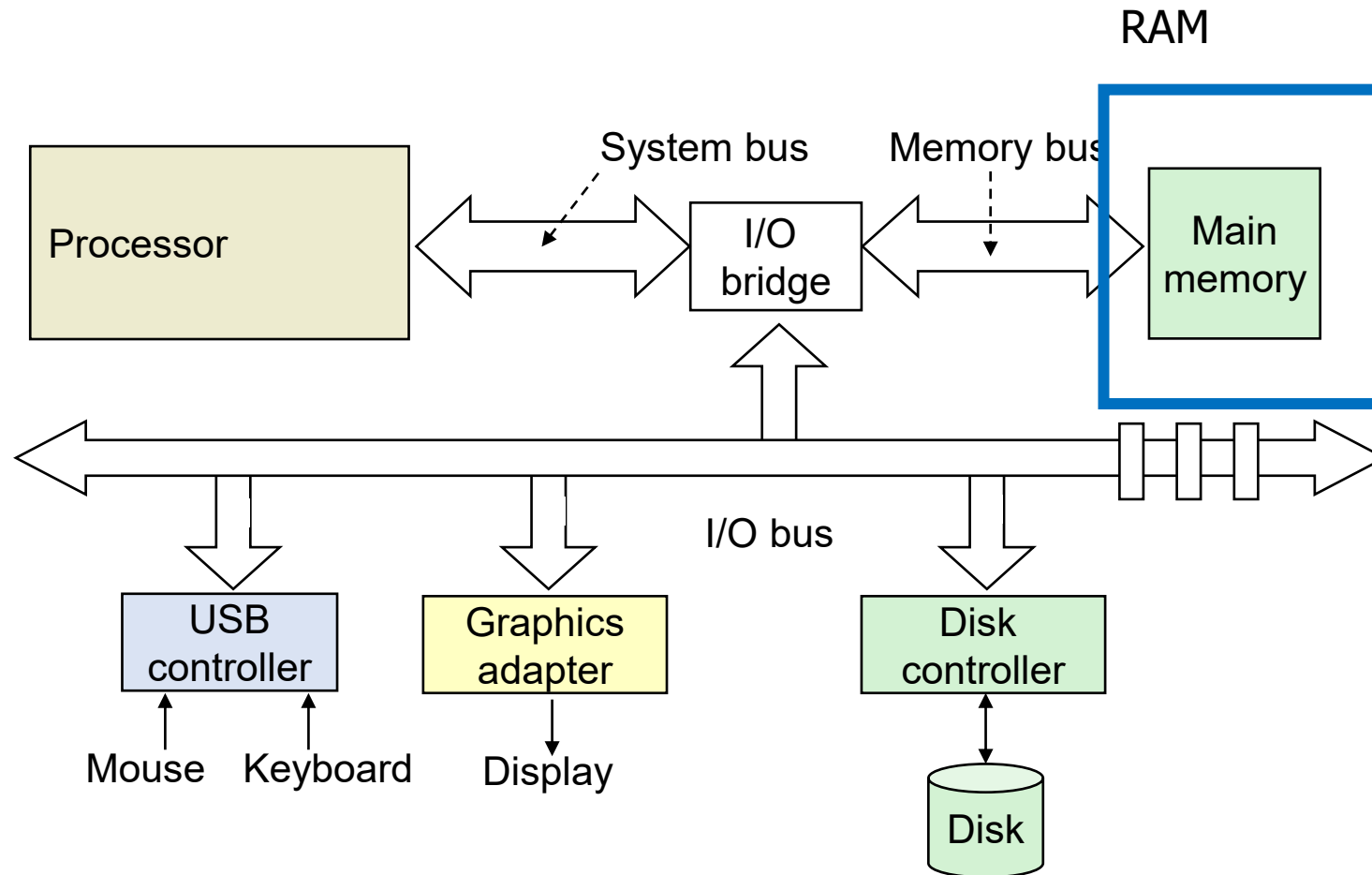
- Fastest memory on computer

- Also most expensive and takes up most space per bit
- Typically used for registers and cache memories

Cache memories (a.k.a. caches)

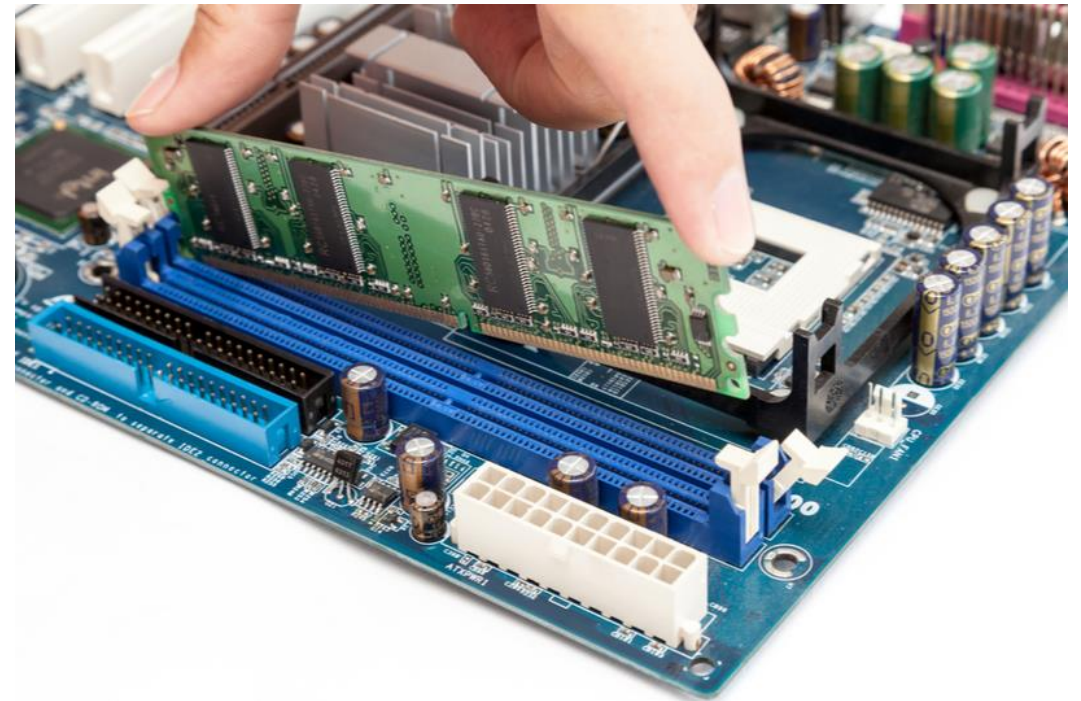
- Memory in the processor that duplicates portions of main memory
 - Goal: speed up accesses to frequently used data
 - Complicated decision: what should be in the cache at any given moment
 - Many configurations to try to improve the “hit rate” on data
- Not general-purpose memory
 - Can't be directly accessed and doesn't have memory addresses
 - Hardware automatically chooses what's in the cache
- Cache sizes
 - Kilobytes to Megabytes of storage: ~1-100 MB total in modern processors

Tour of computer memory



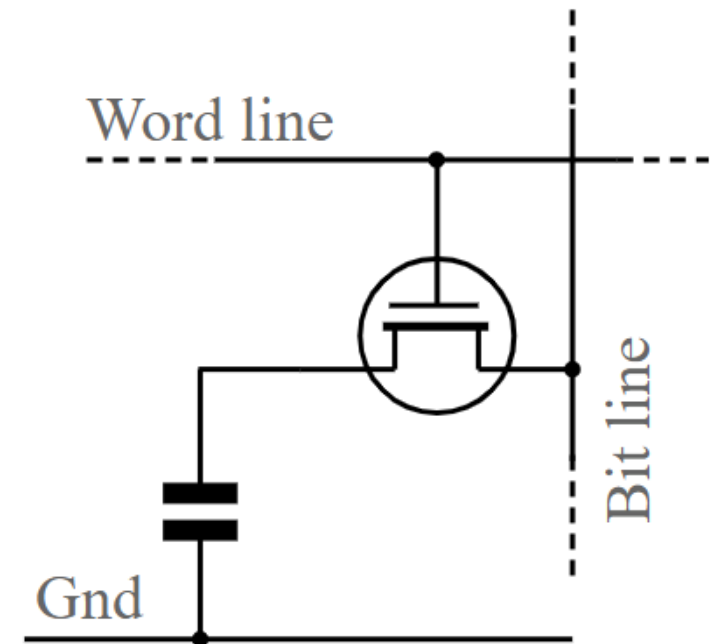
Main memory

- The “RAM” in your computer
 - Random Access Memory
 - Can access any byte you want in “random” order
- Typically measured in GB
 - 1-128 GB
 - Some special purpose systems may have MUCH less or more
- This is the “array of bytes” we’ve been using in assembly
 - Program memory lives in RAM



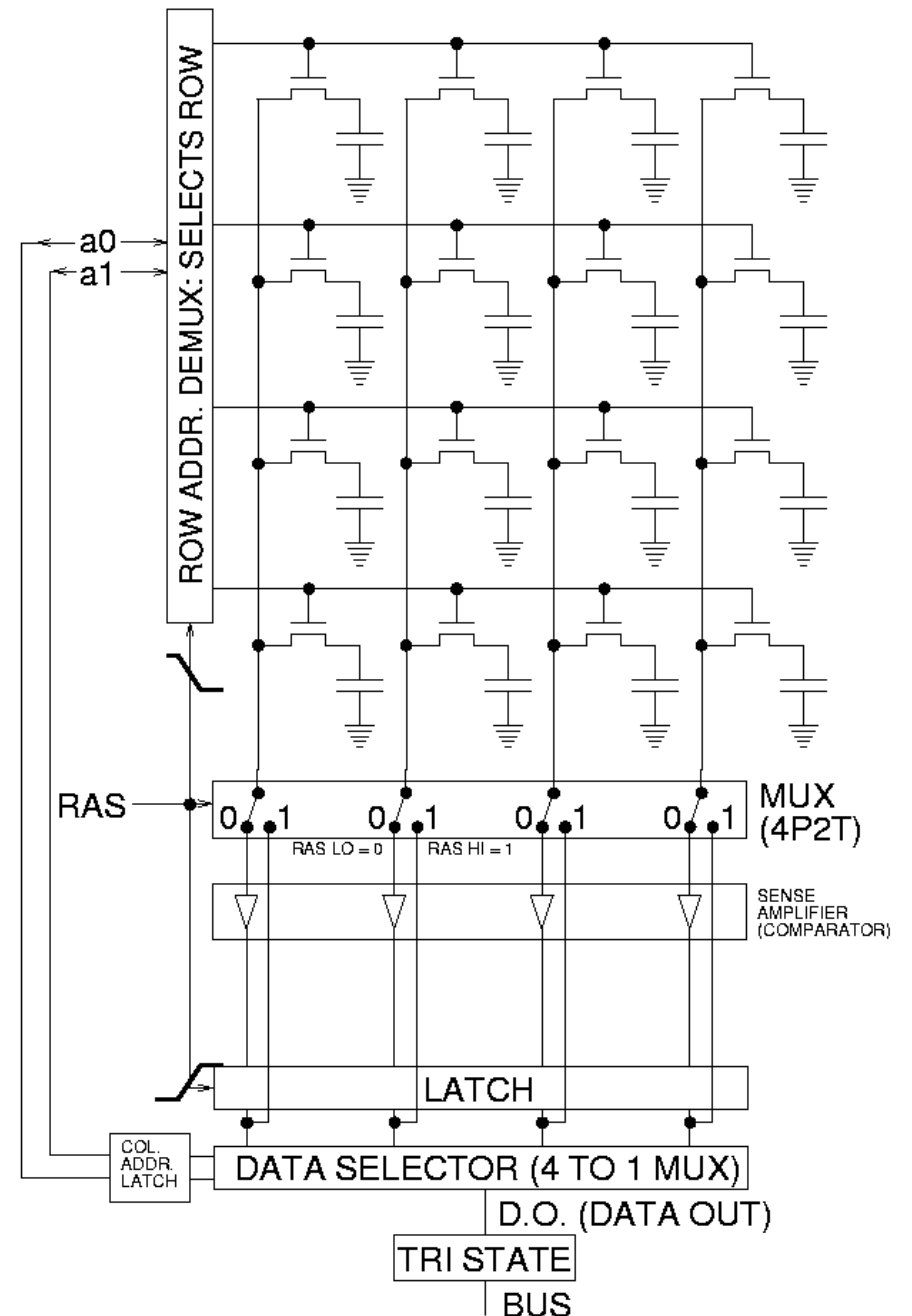
Main memory technology: DRAM

- Dynamic RAM (DRAM)
 - Each cell stores a bit as a charge in a capacitor
 - Capacitors lose charge; each cell must be refreshed every 10-100 ms
 - More sensitive to disturbances (EMI, radiation, ...) than SRAM
- Slower than SRAM, but cheaper and denser
 - ~100x slower than registers
- Typically used for main memory



Accessing DRAM

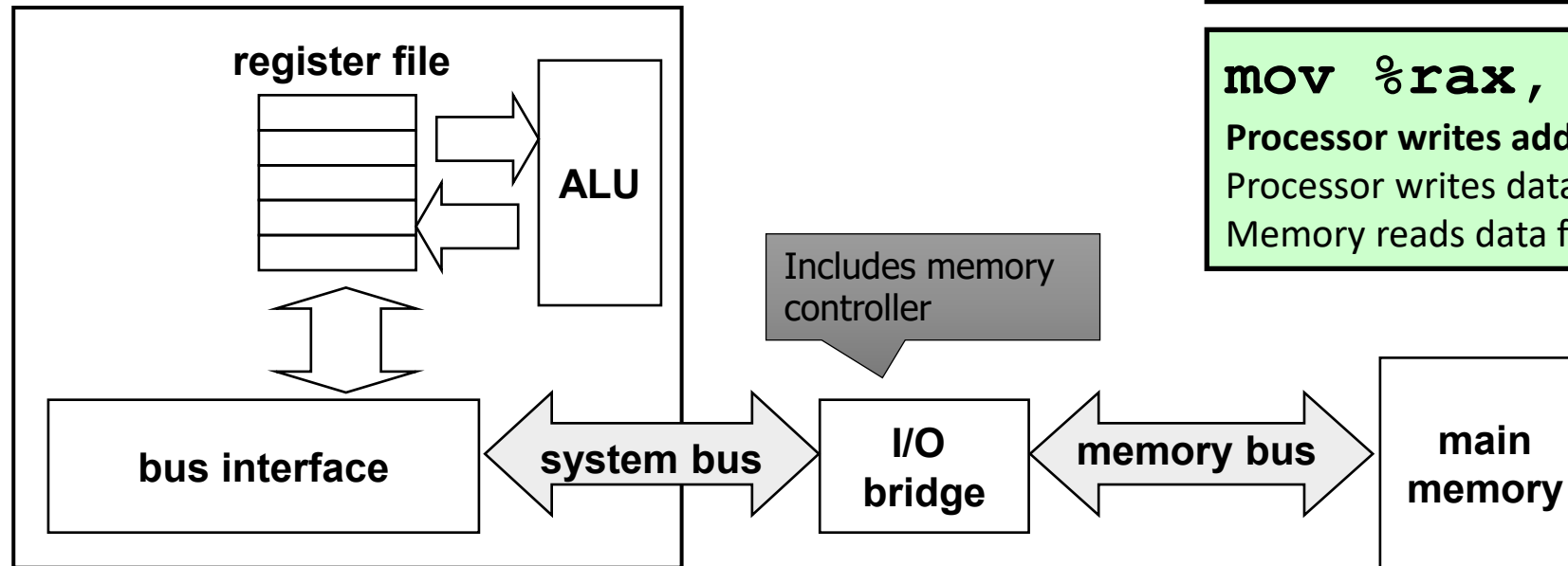
- Read entire row of data at a time
 - Large in practice, kilobytes
- Select actual bytes that are wanted
 - Possibly modifying those bits
- Write row back to memory
 - Must always happen!
 - Reading is destructive



Connecting main memory and the processor

- Data flows between main memory and processor over “buses”
 - A collection of parallel wires that carry address, data, and control signals
 - Typically shared by multiple devices

CPU chip



```
mov (%rdi), %rax {load}
```

Processor writes address A to the bus

Memory writes data to the bus

Processor reads data from the bus

```
mov %rax, (%rdi) {store}
```

Processor writes address A to the bus

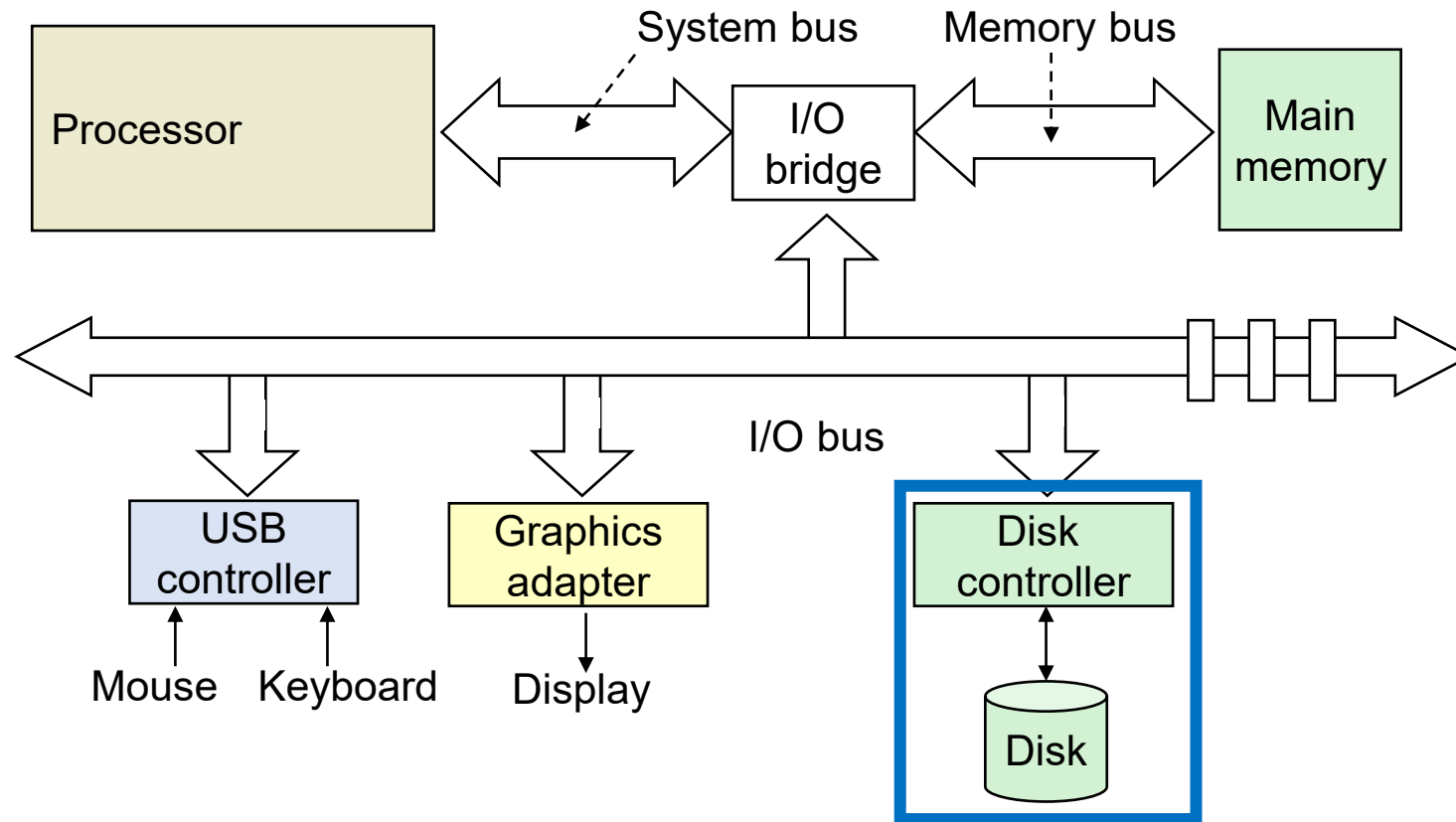
Processor writes data to the bus

Memory reads data from the bus

Sidebar: what's going on with RAM prices?

- Limited supply
 - Fabs make computer chips
 - There really aren't that many of them, and only some of those make RAM
 - ~3 big RAM manufacturers
- Increasing demand
 - Machine Learning needs lots of memory
 - Particularly it needs faster, more expensive memory
- Result: greatly increased costs
 - 32 GB of RAM: 2024 - \$80, 2026 - \$360 🤖

Tour of computer memory



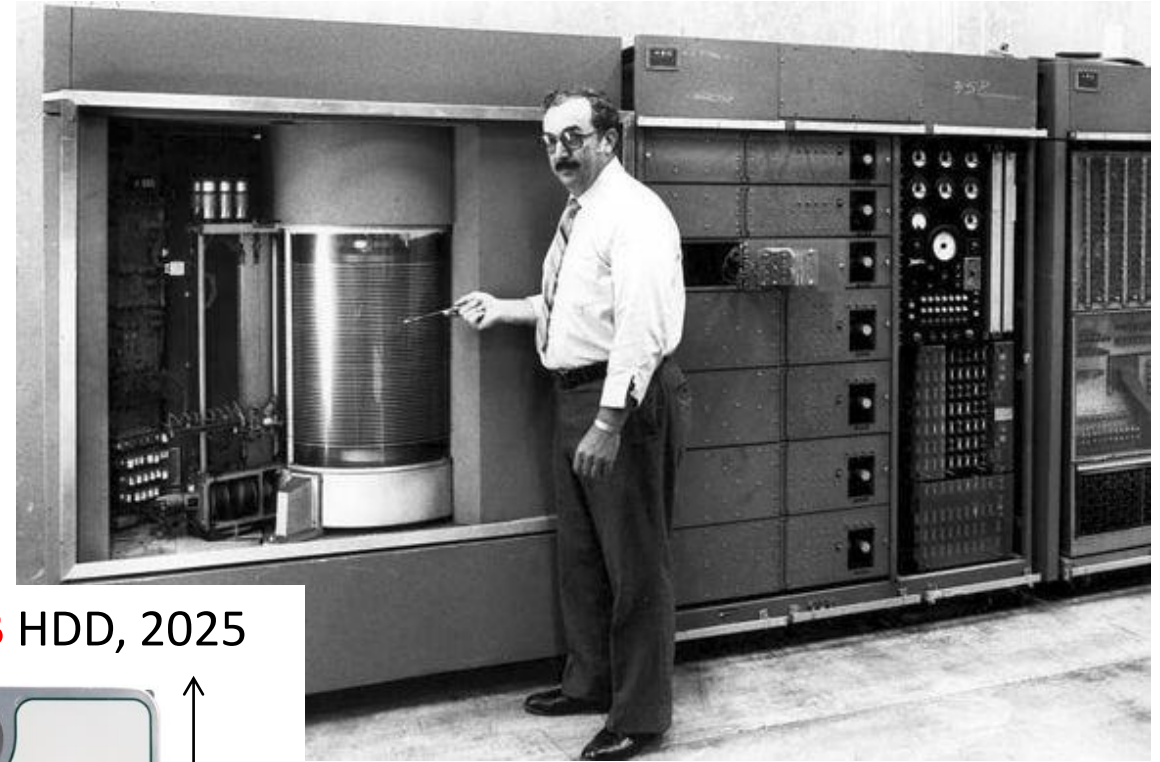
Disk:
Hard drive or SSD

Disk storage

IBM 350 Disk Storage Unit
First disk drive, 1956

Storage?
5MB! (these slides are 8 MB)

- Workhorse storage devices
 - Terabytes in size in modern computers
 - Milliseconds to read
 - 100,000x longer than reading from RAM
 - 10,000,000x longer than reading from registers



Seagate **24TB** HDD, 2025

- Used for filesystem storage
 - Running programs do not directly use disk except when interacting with files

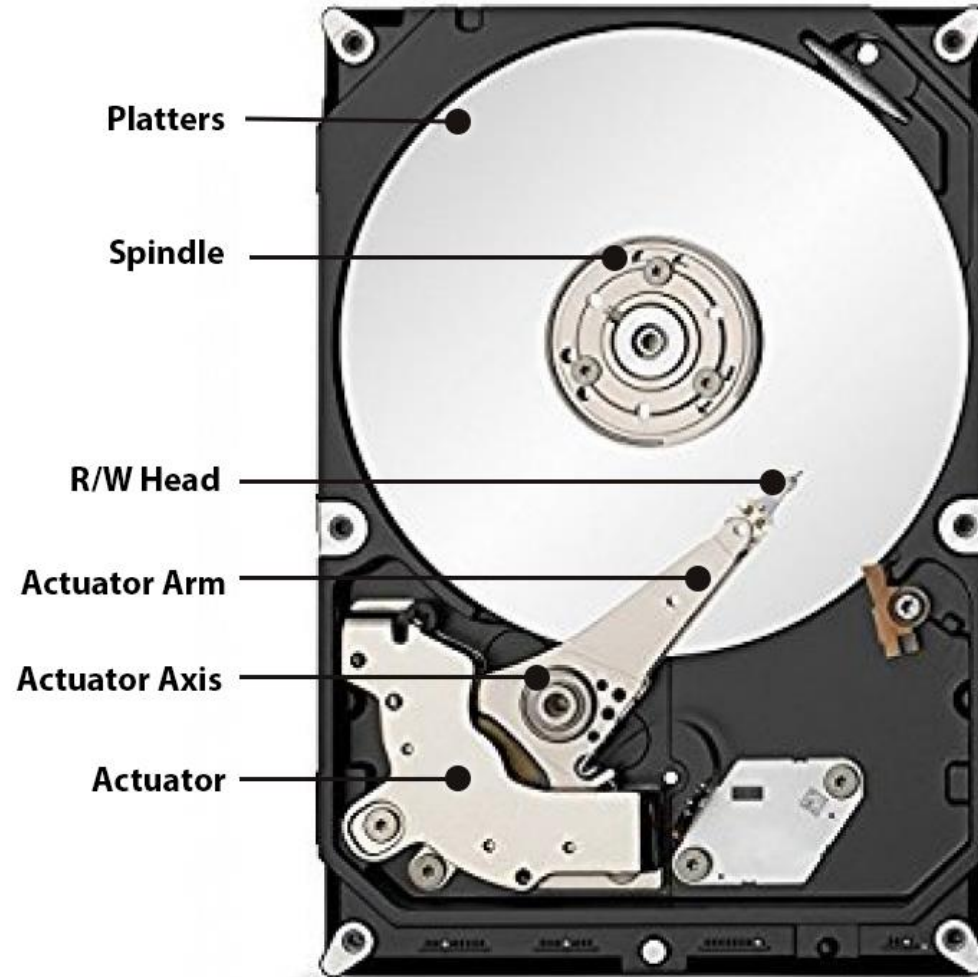


5.75 inches
(146 mm)

Basic mechanisms the same!

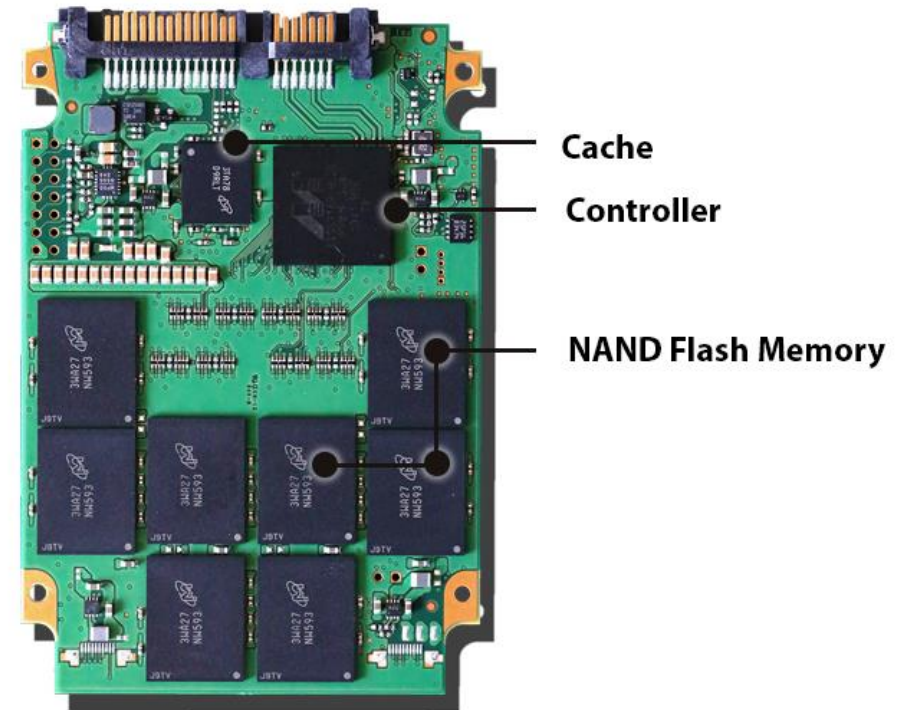
Disk types

HDD 3.5"



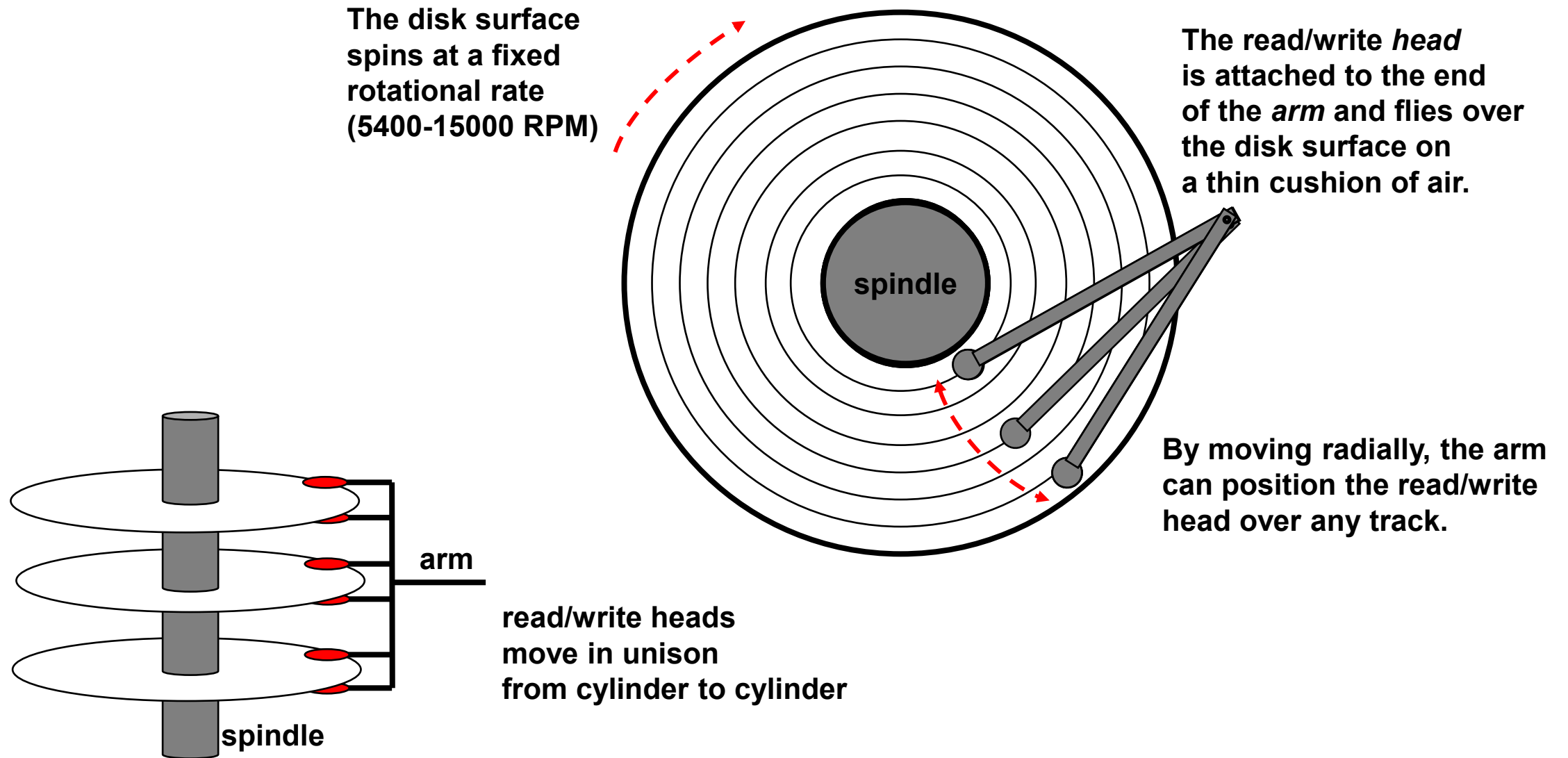
Shock resistant up to 55g (operating)
Shock resistant up to 350g (non-operating)

SSD 2.5"



Shock resistant up to 1500g
(operating and non-operating)

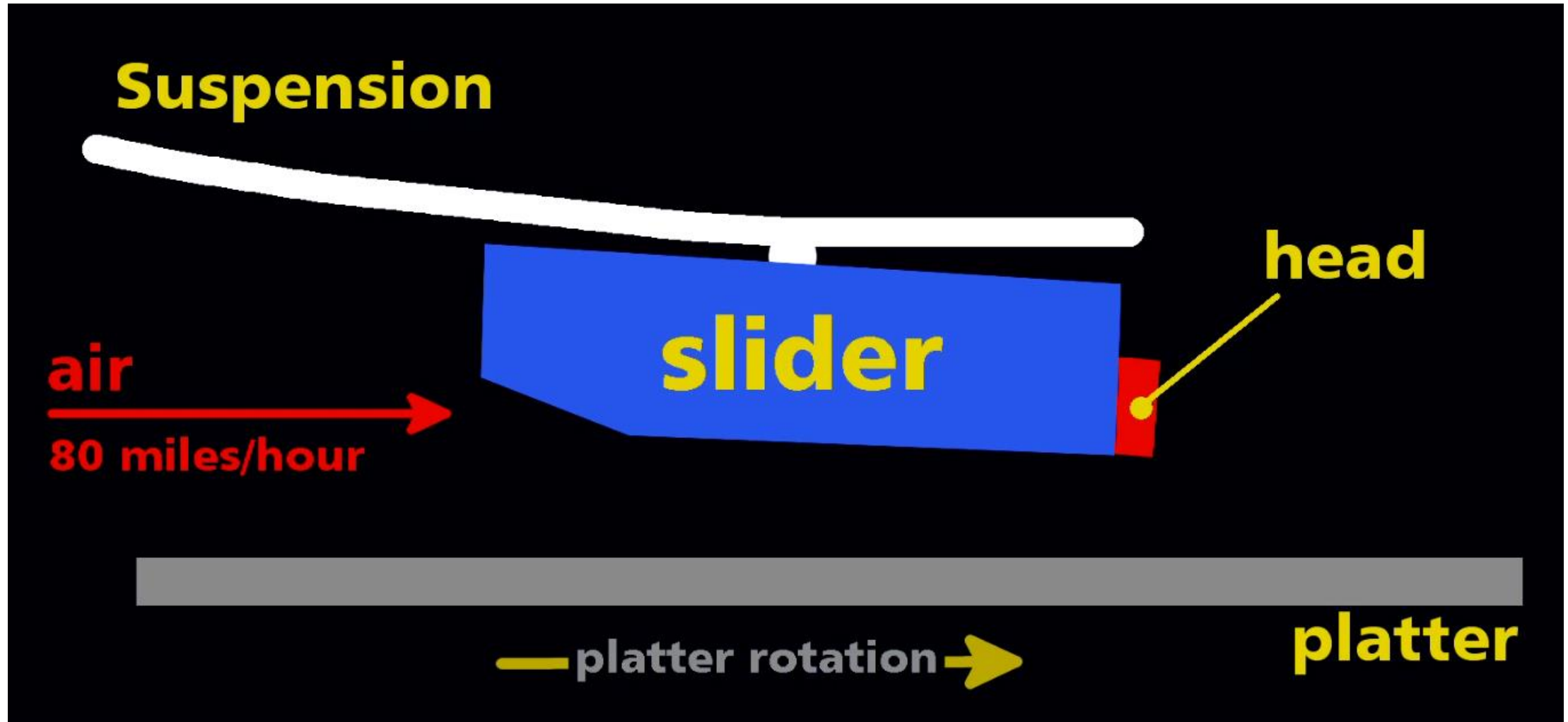
Hard disk drive operation



Animation of Disk Access



Disk heads read/write data to the platters

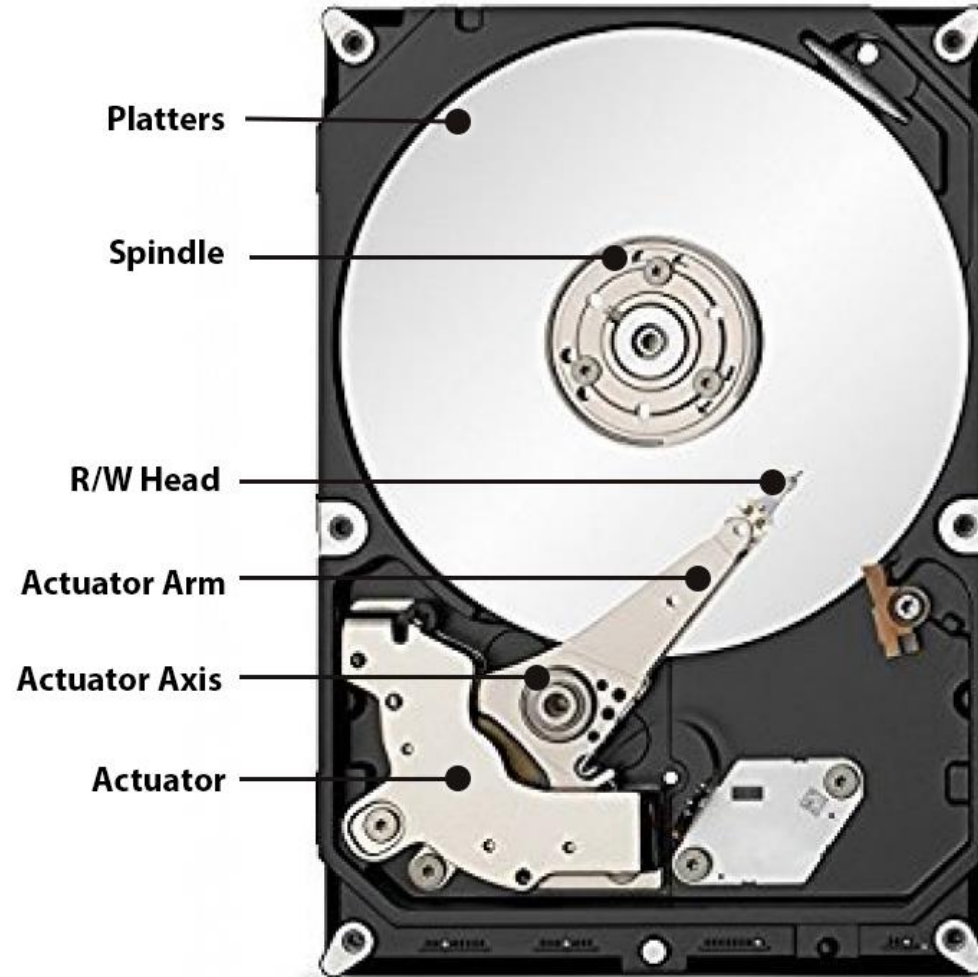


Accessing data on disk

- Three-step process
 - Read head needs to move to the right cylinder: ***seek time***
 - Platter turns until start of sector under the read head: ***rotational latency***
 - Sector passes under read head and data is read: ***transfer time***
- Time cost dominated by seek time and rotational latency
 - **Consequence:** reading first bit of a sector is expensive
 - But reading the rest of the sector is basically free!
- When using disks, best to favor large sequential reads/writes
 - Terrible for random access! (reading a little bit here, a little bit there)
 - Fits file access well (read/write in order)
 - But overall still really slow compared to main memory or registers

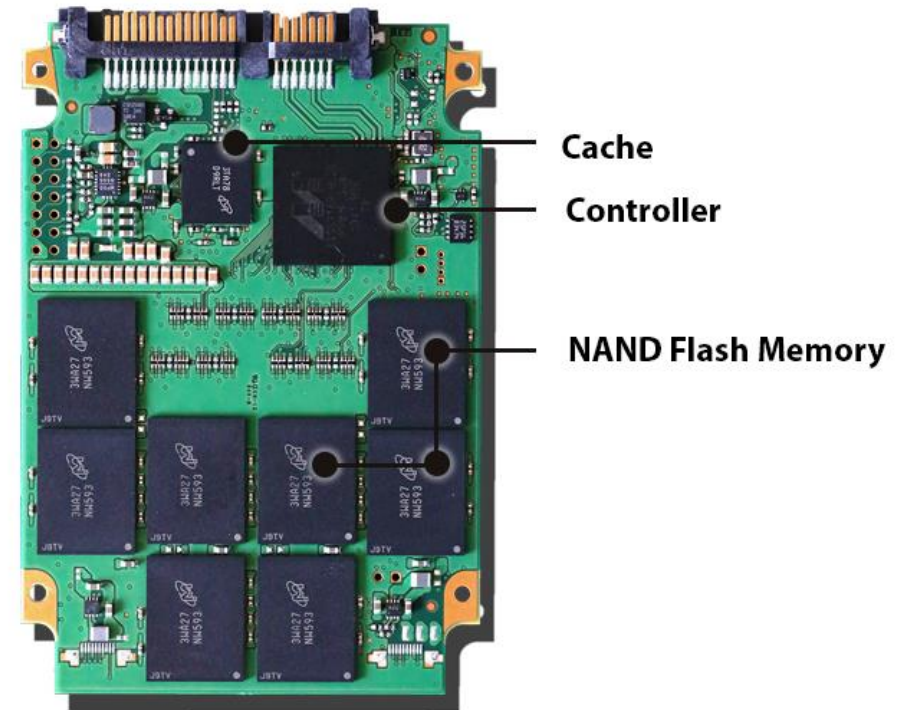
Disk types

HDD 3.5"



Shock resistant up to 55g (operating)
Shock resistant up to 350g (non-operating)

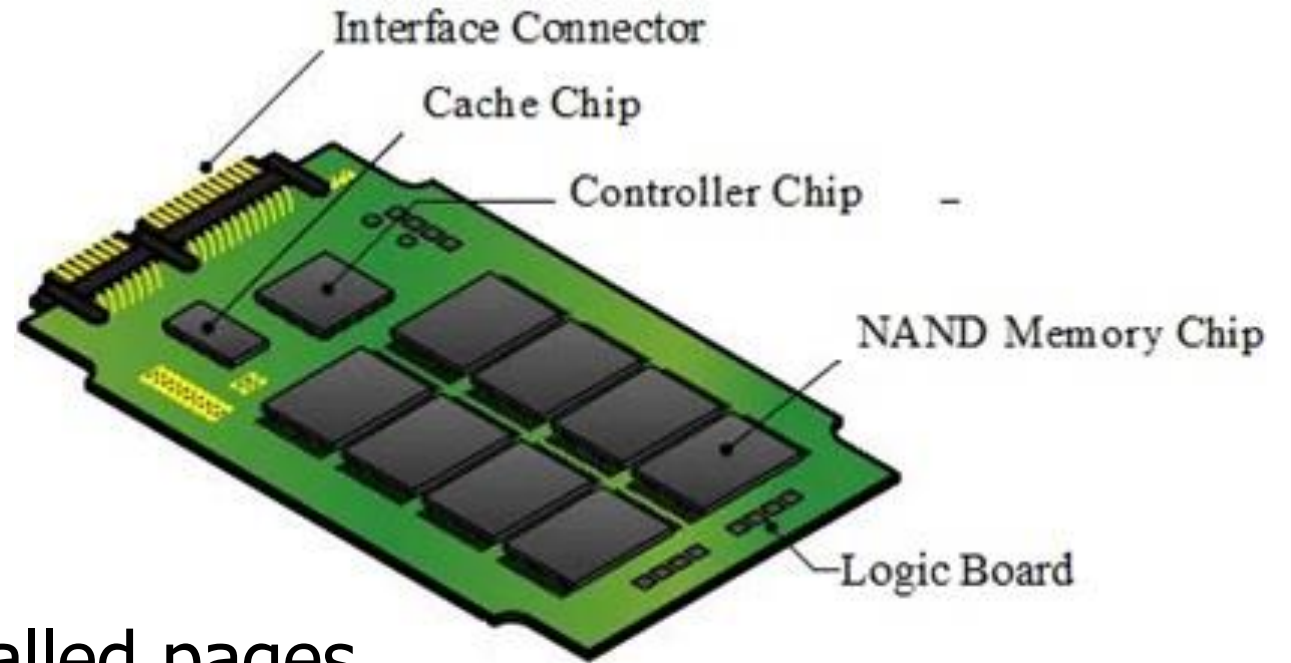
SSD 2.5"



Shock resistant up to 1500g
(operating and non-operating)

Solid State Drives (SSDs)

- Flash memory
 - Just like Flash Drives
- Stores data in chips, so it's much faster to access
- Organize data into chunks, called pages
 - Each page is 4-16 KB in size
 - Entire page must be rewritten at once (not individual bytes)
- Pages wear out after repeated writes (~100K writes total) and can no longer be used



SSDs vs Rotating Disks

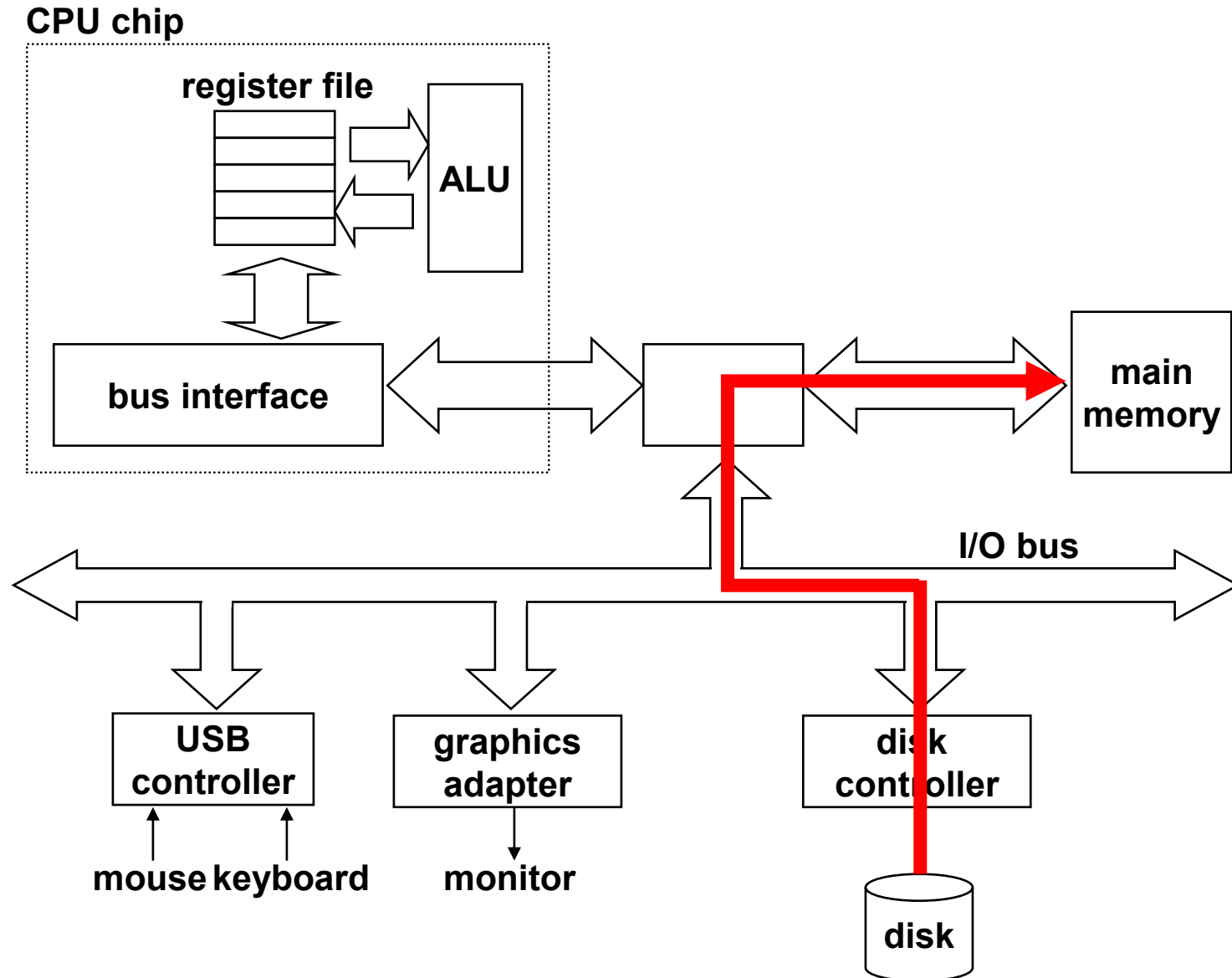
- Advantages
 - No moving parts → faster, less power, more rugged
- Disadvantages
 - Have the potential to wear out
 - Mitigated by “wear leveling logic” in flash translation layer
 - More expensive per byte (but getting cheaper)
 - 2026: HDD \$0.016 per GB, SSD \$0.070 per GB (about 4x the cost)
 - (also roughly 1.5x the cost from last year when I made this slide)
- Applications
 - Portable electronics (phones, tablets, etc.)
 - Began to appear in desktops and servers circa 2007
 - Now common on laptops as well

Biggest speed improvement
to your computer:

- Get an SSD

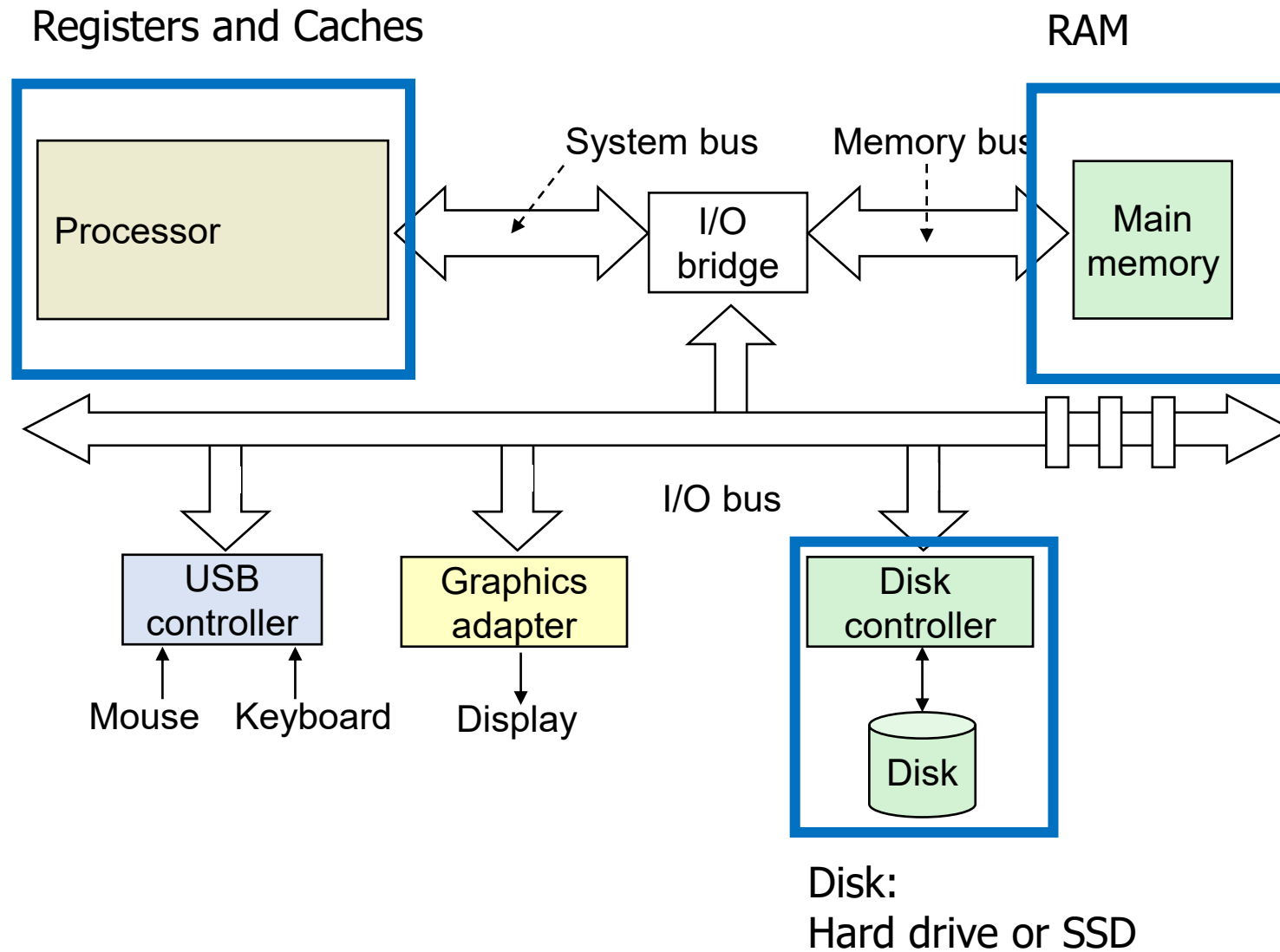
<https://diskprices.com/>

Reading memory from disk



- Data from disk is always read into Main Memory
- Direct Memory Access (DMA)
 - Processor starts a read and then returns to programs
 - Disk performs the read, transfers data, then notifies processor when done

Tour of computer memory



Break + Question

- How do you make an SSD with a longer lifetime (more writes)?
 - Without changing any of the physics of how it works

Break + Question

- How do you make an SSD with a longer lifetime (more writes)?
 - Without changing any of the physics of how it works
- Secretly make it larger than it claims to be
 - e.g. 200 GB when it claims to be 100 GB
- Behind the scenes move around memory as necessary so the device can still hold 100 GB even if half of the flash is
 - Maintain a mapping of which data is located where

Outline

- What's in a processor? Assembly-to-Transistors
- Memory Technologies
- **Memory Hierarchy**
- Caching Concept

Computing timescales

- Assuming 4 GHz processor, **Instruction (with registers): 0.25 ns**

**Jeff Dean
(Google AI):
"Numbers Everyone
Should Know"**

Reminder:
1,000,000,000 ns per second

L1 cache reference	0.5 ns
Branch mispredict	5 ns
L2 cache reference	7 ns
Mutex lock/unlock	25 ns
Main memory reference	100 ns
Compress 1K bytes with Zippy	3,000 ns
Send 2K bytes over 1 Gbps network	20,000 ns
Read 1 MB sequentially from memory	250,000 ns
Round trip within same datacenter	500,000 ns
Disk seek	10,000,000 ns
Read 1 MB sequentially from disk	20,000,000 ns
Send packet CA->Netherlands->CA	150,000,000 ns

Jim Gray's storage latency analogy

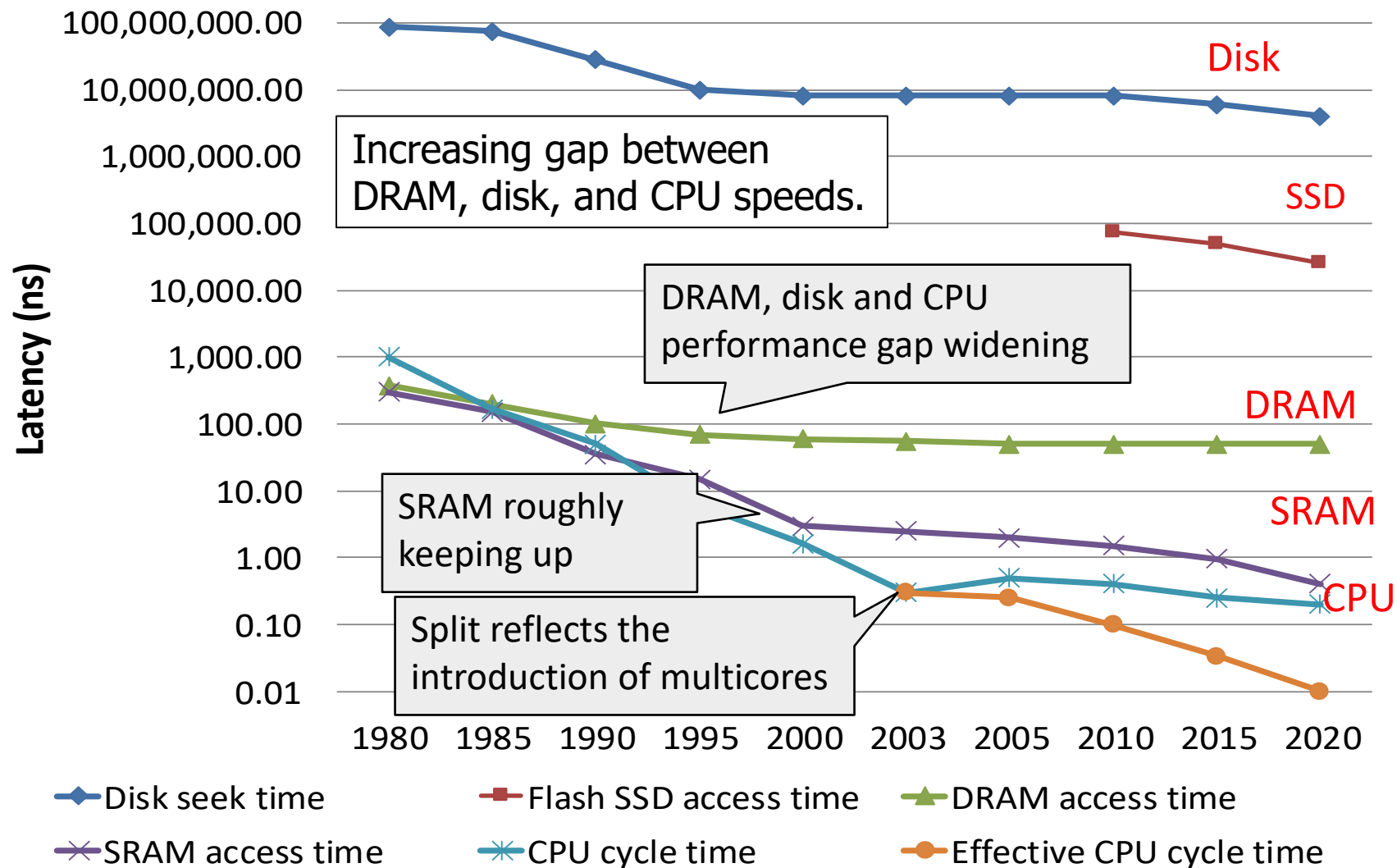
- How “far” is data for the CPU, converted to human scale

Storage	Distance	Time
Registers	In my apartment	~1 minute
On-chip cache	Across the street	2-4 minutes
On-board cache	A few blocks away	10 minutes
Main memory	In Milwaukee	1.5 hours
Disk	On Mars	2 years
Tape	On Kepler-76b	2000 years (at speed of light)

Jim Gray
Turing Award 1998



The CPU-Memory gap



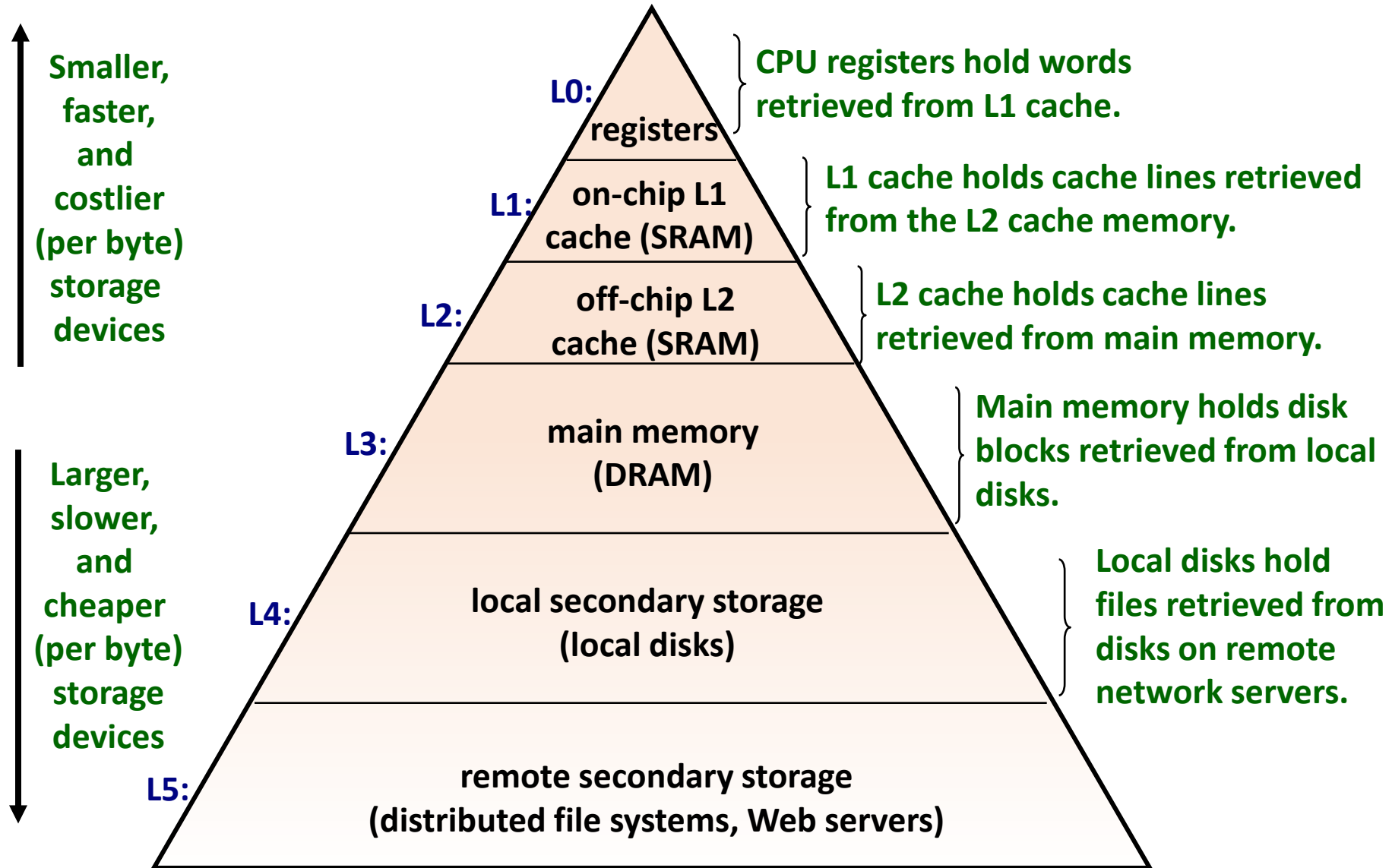
The CPU-Memory Gap

- CPUs outspeed memory
 - But they can't compute on data they don't have!
 - If the CPU has to wait for data to reach it, it just sits idle!
 - All that engineering effort, just sitting around waiting...
- Challenge: get data to the CPU despite "slow" memory
 - So the CPU can work at full speed, without waiting for data
- Two-pronged strategy
 - ***Memory hierarchy***: keep data we need closer to the CPU
 - ***Locality of reference***: predict which data we're likely to need

Memory hierarchy

- Some fundamental and enduring properties of systems
 - The faster the storage, the more expensive (\$) it is
 - The faster the storage, the smaller (capacity) it is
 - The gap between processor and main memory speed is widening
- Key idea: keep the data you need the most in fast memory!
 - Data you only need from time to time can be in slow memory, no big deal
 - Most used data goes in registers
 - Least used data goes to disk
- Analogy: kitchen ingredients I use
 - Salt, all the time: it sits out on the counter
 - Oregano, frequently: front of the cabinet
 - Turmeric, occasionally: back of the cabinet
 - Brown sugar, sometimes: somewhere in the pantry
 - Saffron, never: I can go buy some if I do need it

Memory hierarchy



Outline

- What's in a processor? Assembly-to-Transistors
- Memory Technologies
- Memory Hierarchy
- **Caching Concept**

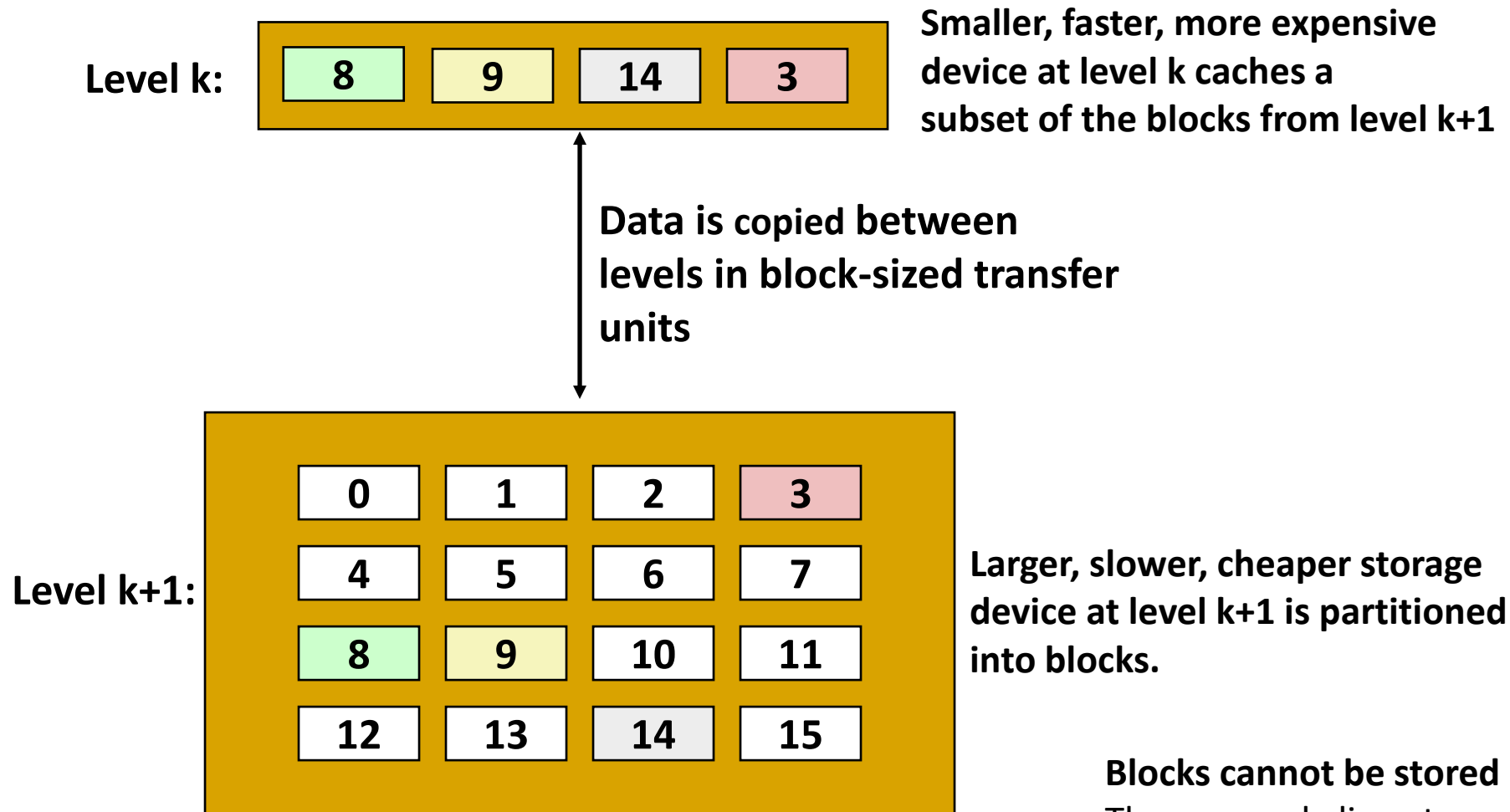
Caching, general principle

- **Cache:** A smaller, faster storage device that acts as a staging area for a subset of the data in a larger, slower device
 - Data lives in both places (typically)
 - When the consumer (e.g., CPU) reads the data, it gets it from the smaller, faster storage
 - If the data we want is not in the cache, we pay the full cost of bringing it over from the larger, slower storage into the smaller, faster storage
 - The hope: we don't need to do it too often
- Fundamental idea in systems. Shows up all over!
 - Memory hierarchies
 - Content delivery networks (CDNs) on the Internet (Akamai, Cloudflare, etc.)

Caching in a memory hierarchy

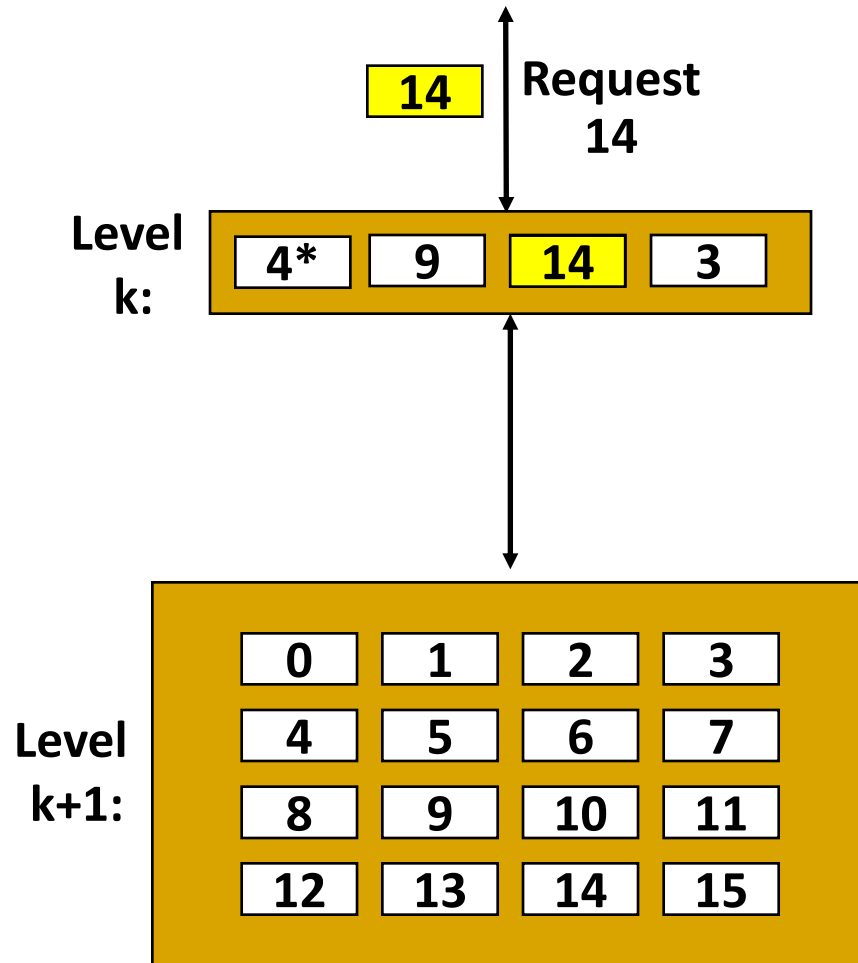
- Fundamental idea of a memory hierarchy
 - For each **k**, the faster, smaller device at level **k** serves as a cache for the larger, slower device at level **k+1**
 - L2 cache memory as a cache for main memory
 - Main memory as a cache for disk, etc.
 - Each level stores some of the most frequently accessed data
 - The closer the cache is to the processor, the “hotter” the cached data

Caching in a memory hierarchy



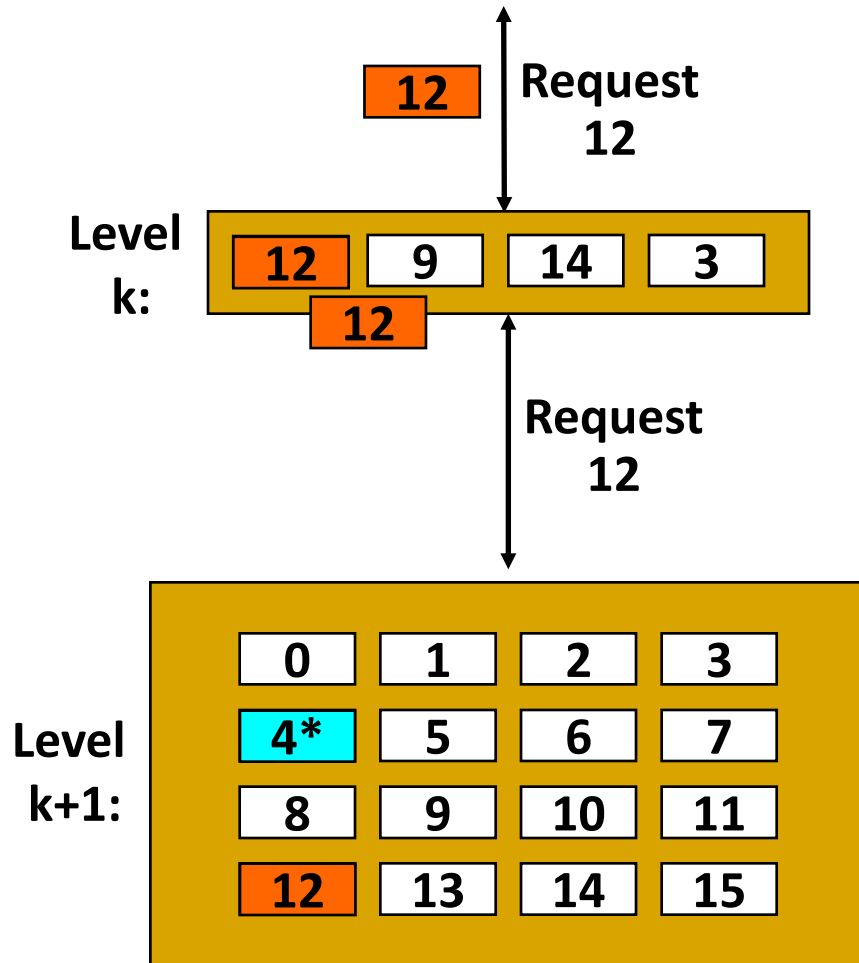
Blocks cannot be stored in an arbitrary location!
They can only live at one of a fixed set of locations.
In this example: they must be in the same “column” for both levels.

General caching concepts



- Program needs object **d**, which is stored in some block **b**
- **Cache hit**
 - Program finds **b** in the cache at level **k**
e.g., block 14

General caching concepts



"*" means the block is *dirty*
(i.e., it has been modified)

- Program needs object **d**, which is stored in some block **b**
- **Cache hit**
 - Program finds **b** in the cache at level **k**
e.g., block 14
- **Cache miss**
 - **b** is not at level **k**, so the level **k** cache must fetch it from level **k+1**,
e.g., block 12
 - If the level-k cache is full, then some current block must be replaced (**evicted**). Which one is the "victim"?
 - Here, we pick 4; same column as 12
 - 4 is "dirty", need to write back to k+1
 - More on this next lecture

Memory hierarchies make memory fast and large

- Why do memory hierarchies work?
 - Programs tend to access the data at level **k** more often than they access the data at level **k+1**
 - Thus, the storage at level **k+1** can be slower, and thus larger and cheaper per bit
- Net effect:
 - A large pool of memory that costs as little as the cheap storage near the bottom, but that serves data to programs at the rate of the fast storage near the top
 - Best of both worlds!

What causes a cache miss?

- **Cold (compulsory) miss**

- Cold misses occur when a block is accessed for the first time
- No one ever accessed it, so there wasn't any reason to bring it into cache

- **Capacity miss**

- Occurs when the set of active cache blocks (*working set*) is larger than the cache
- There's no way the working set can all fit in the cache, so there will be misses

- **Conflict miss**

- In most caches, blocks cannot be stored in any available slot
- If two blocks need to go in the same slot, need to evict the old one to store the new!
- If after that, we need to access the old block, conflict miss!
 - We had a conflict, evicted a block, and now we miss trying to access that block
- **Note:** can happen even when there is "room" elsewhere in the cache!
- We'll show examples of this next lecture

Break + Question

Miss types: Cold, Capacity, Conflict

- When you first start up a program, it runs really slowly for a few seconds. What kind of cache misses are occurring?
- When you have too many browser tabs open and active, all of the tabs run more slowly. What kind of cache misses are occurring?

Break + Question

Miss types: Cold, Capacity, Conflict

- When you first start up a program, it runs really slowly for a few seconds. What kind of cache misses are occurring?
 - Cold (aka Compulsory) misses. The data has never been loaded before!
- When you have too many browser tabs open and active, all of the tabs run more slowly. What kind of cache misses are occurring?
 - Capacity misses. You have too much data to fit it all in the cache
 - Could be Conflict misses as well, but probably not

Outline

- What's in a processor? Assembly-to-Transistors
- Memory Technologies
- Memory Hierarchy
- Caching Concept