

CS 211: Intro to Computer Architecture

13.2: Memory Access Patterns and Caching

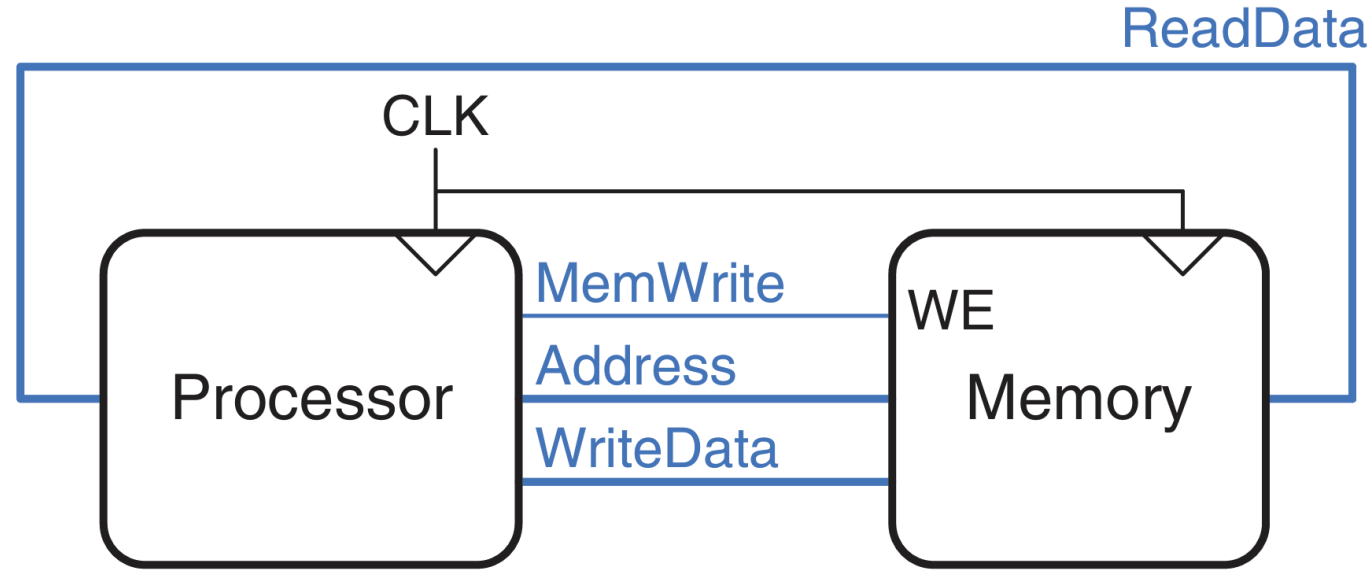
Minesh Patel

Spring 2025 – Thursday 24 April

Announcements

- Ongoing
 - **WA9:** due **Friday (April 18) @ 11:59 pm**
 - **PA5:** due **Monday (May 5) @ 11:59 pm**
- Upcoming
 - **WA10:** TBA: planned for Friday
 - **Final Exam:** May 14th

Our Mental Model (So Far)



- In reality, we have **different types of memory**
- **Today: registers** vs. **caches** vs. **main memory**

Agenda

- **Why Is Main Memory Slow?**
- Caching and the Memory Hierarchy
- Direct-Mapped Caches
 - Where to Place Data
 - How to Query the Cache

Semiconductor Memories

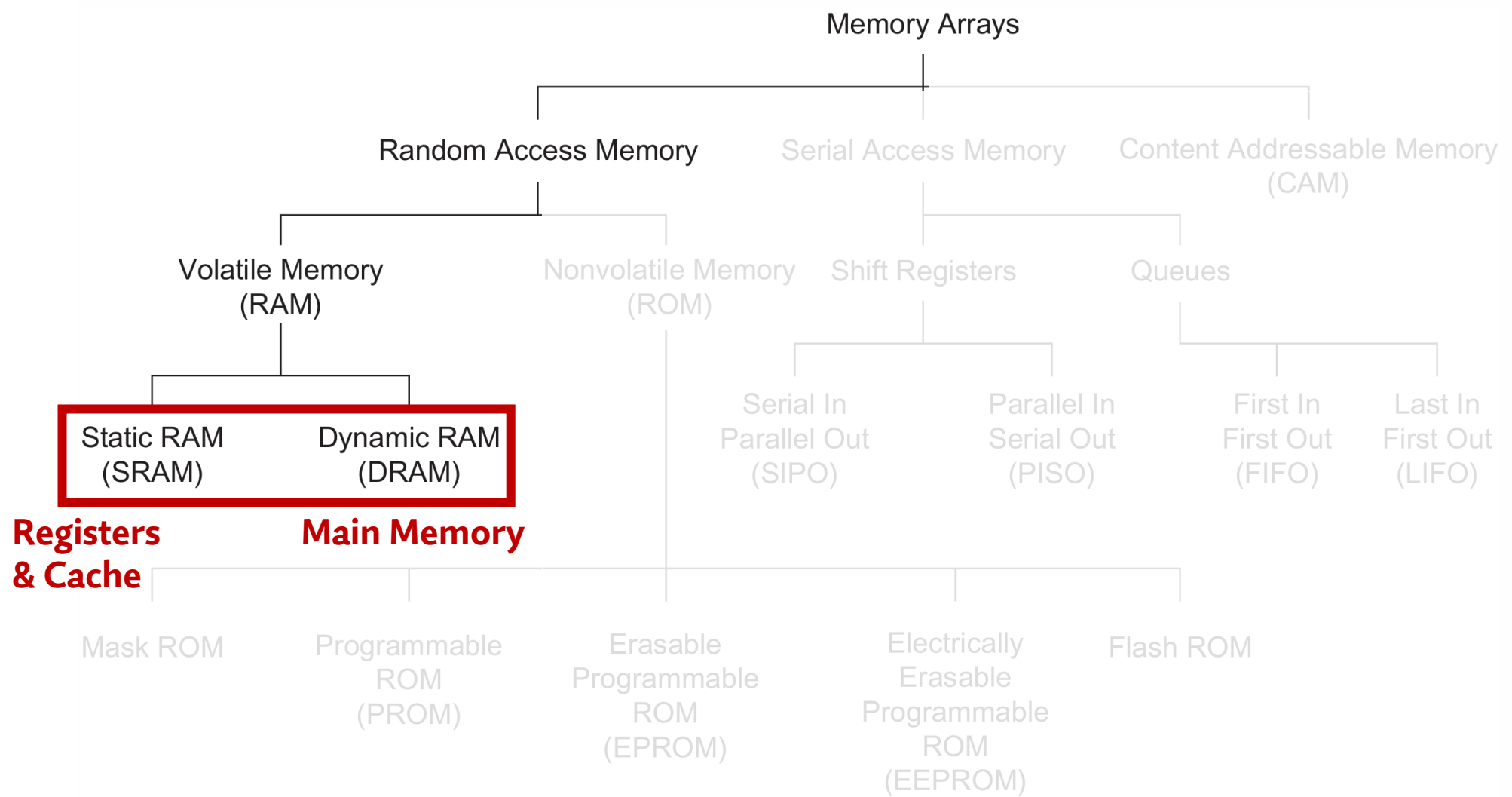
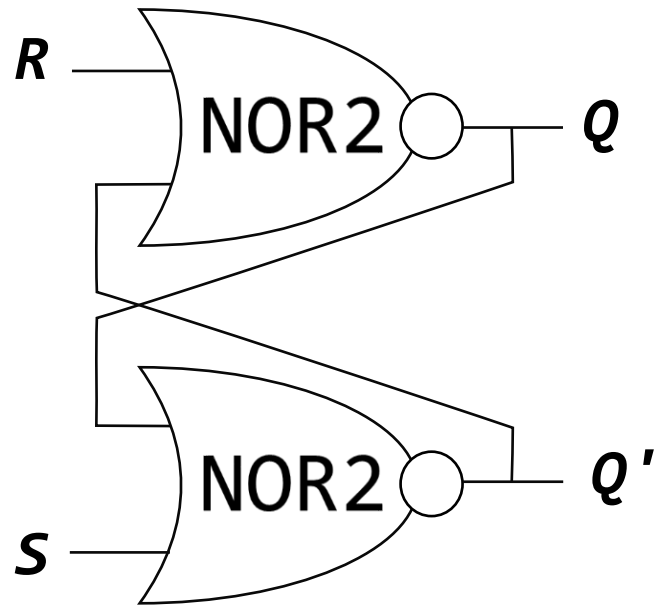


FIGURE 12.1 Categories of memory arrays

Main Memory vs. Registers

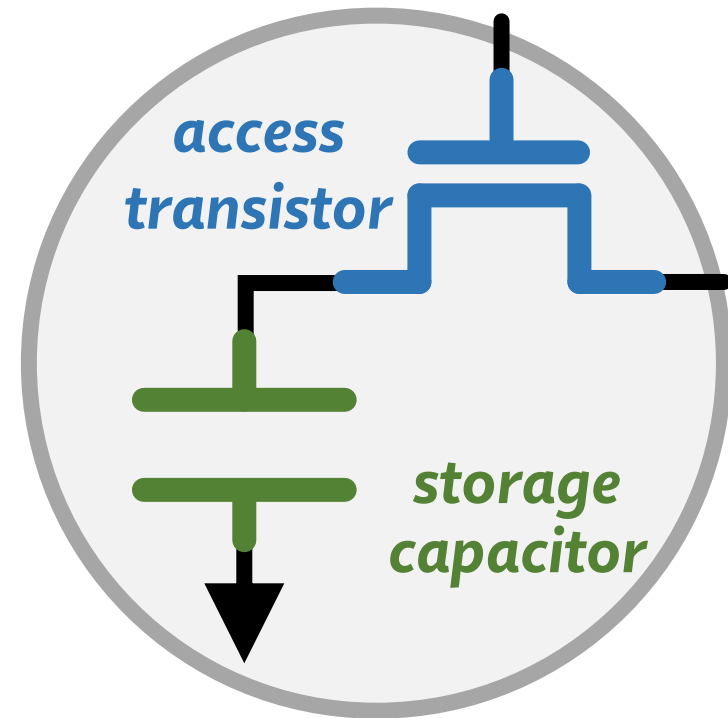
Registers (Flip-Flops)

Built from digital logic



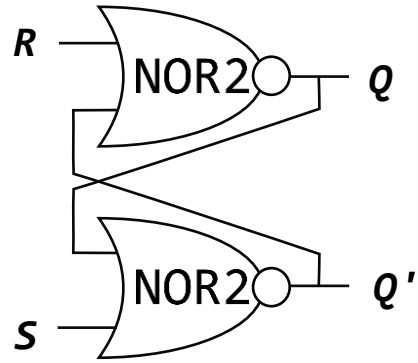
Main Memory (DRAM)

Built from analog circuits

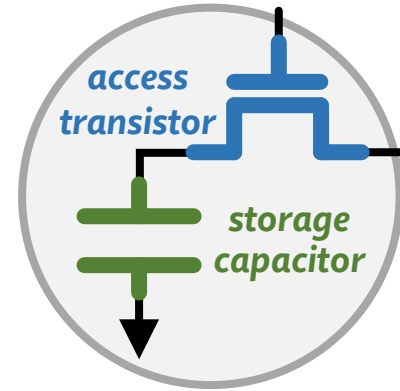


each stores one bit of data

Main Memory vs. Registers



Registers (Flip-Flops)
Built from digital logic



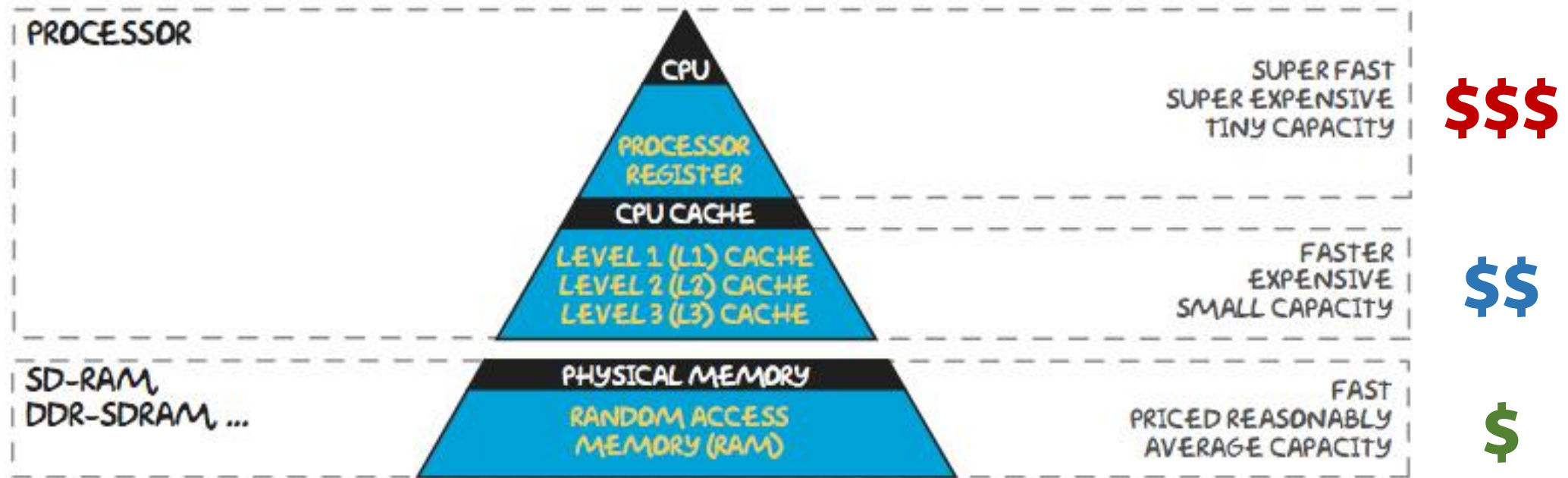
Main Memory (DRAM)
Built from analog circuits

Cost	Expensive
Speed	Fast
Physical Size	Big

Cheap
Slow
Small

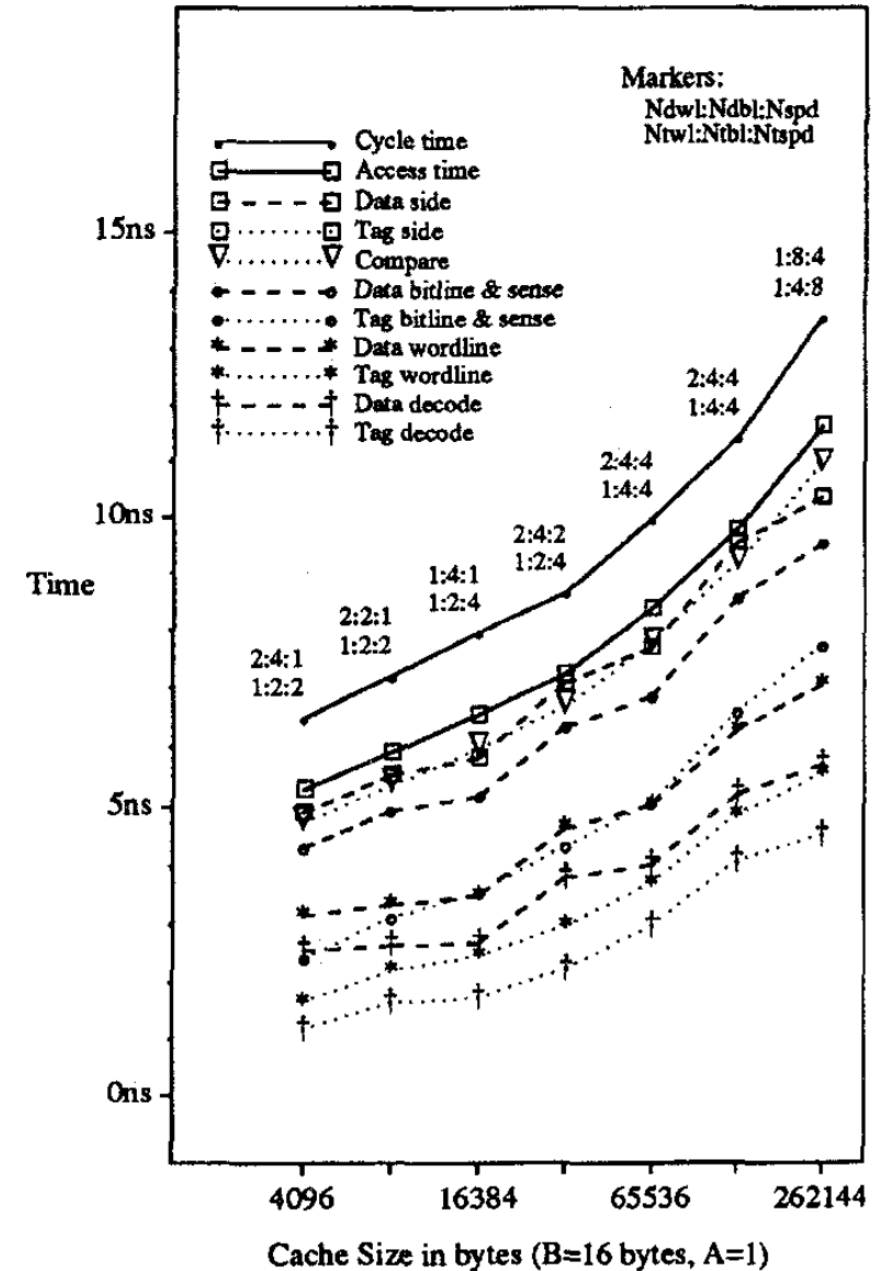
Even If I Were a Billionaire....

- I **still couldn't afford** to keep all data in CPU registers



Beyond Cost: Access Latency

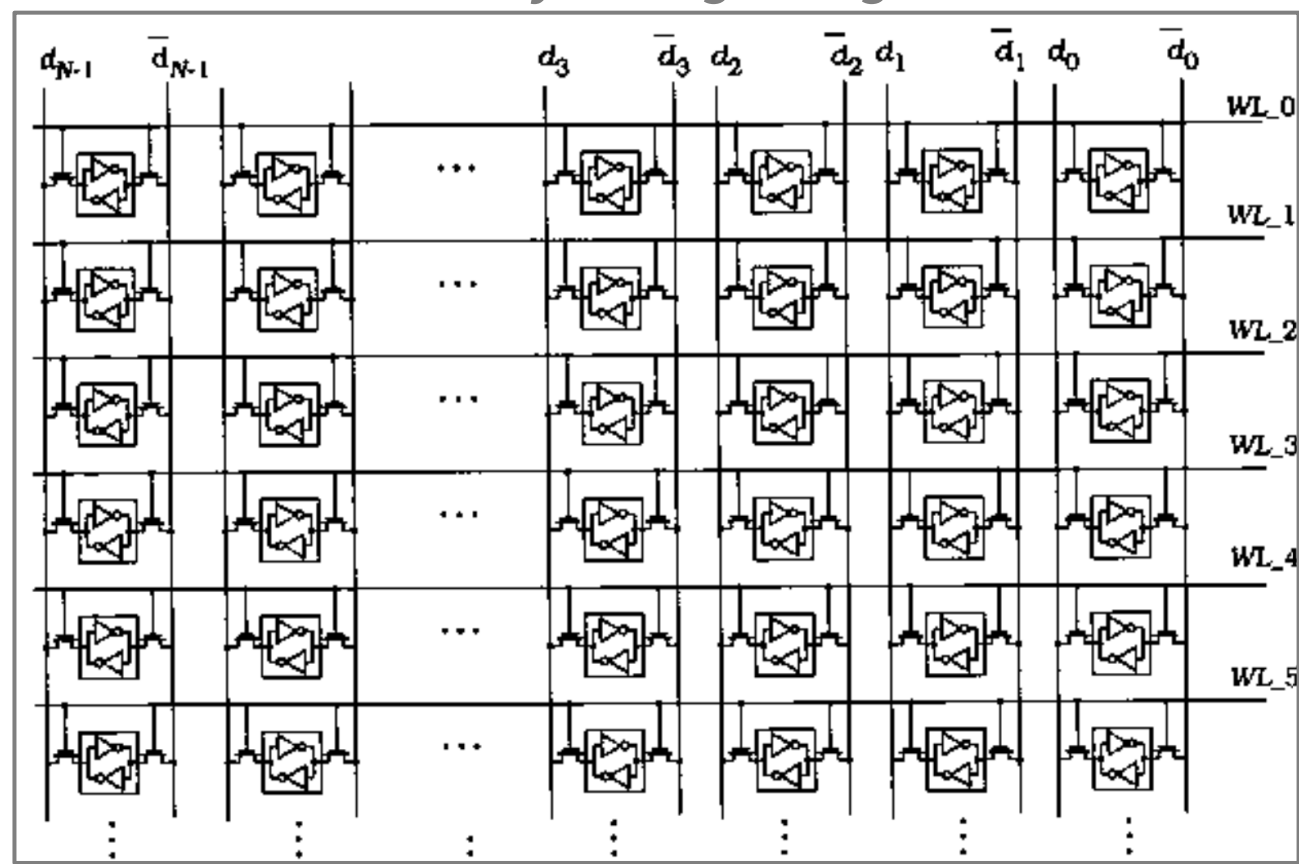
- Fundamentally: **bigger = slower**
 - Electrical signals travel further
 - More wire resistance
 - ...
- **Low-capacity** memories are fastest



Beyond Cost: Density

Static RAM (Registers, Caches)

Built from digital logic



Dynamic RAM (Main Memory)

Built from analog circuits

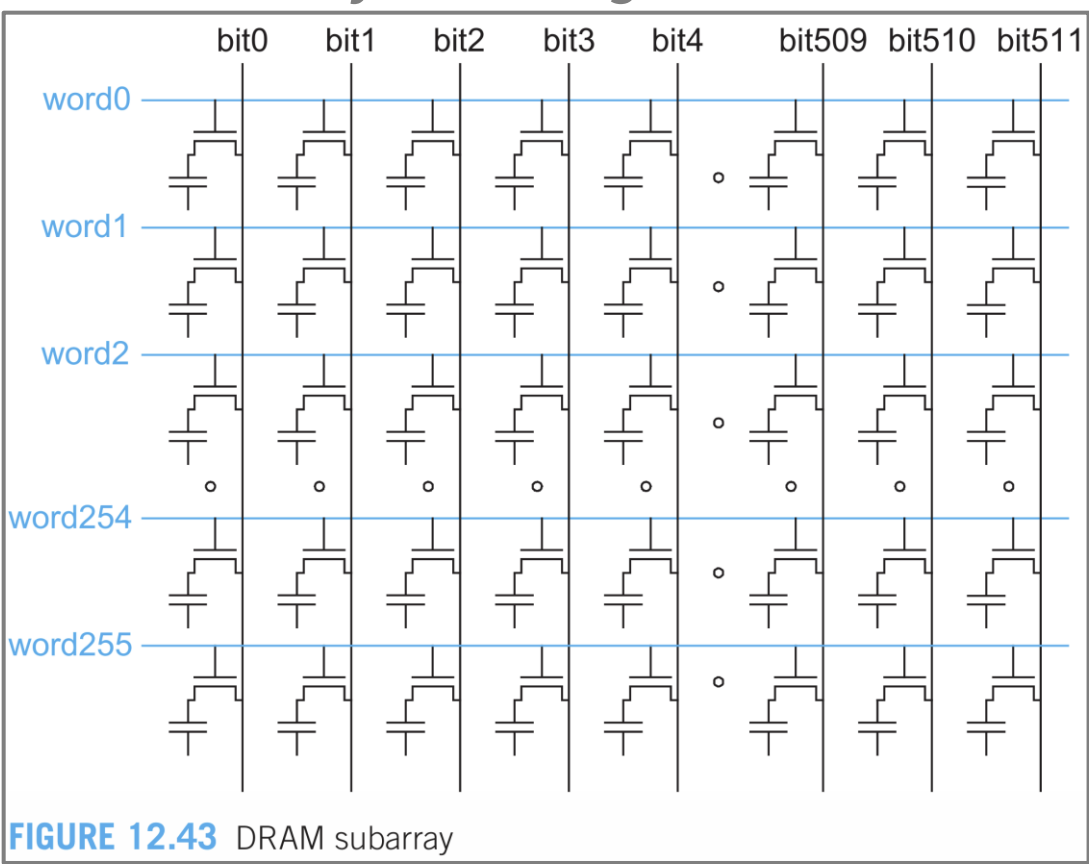
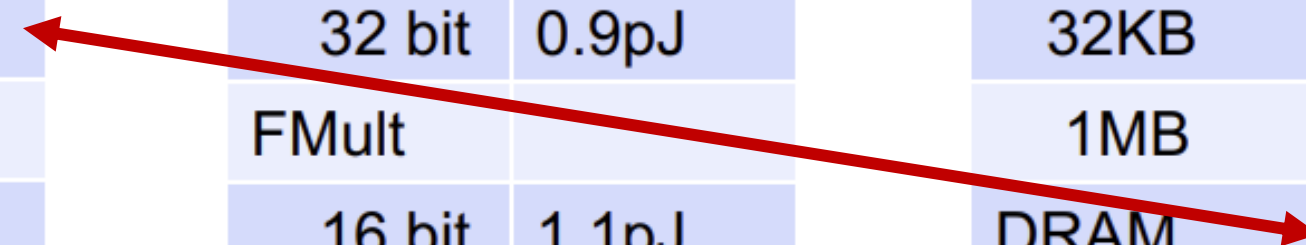


FIGURE 12.43 DRAM subarray

Beyond Cost: Energy

Energy Cost of Various CPU Operations

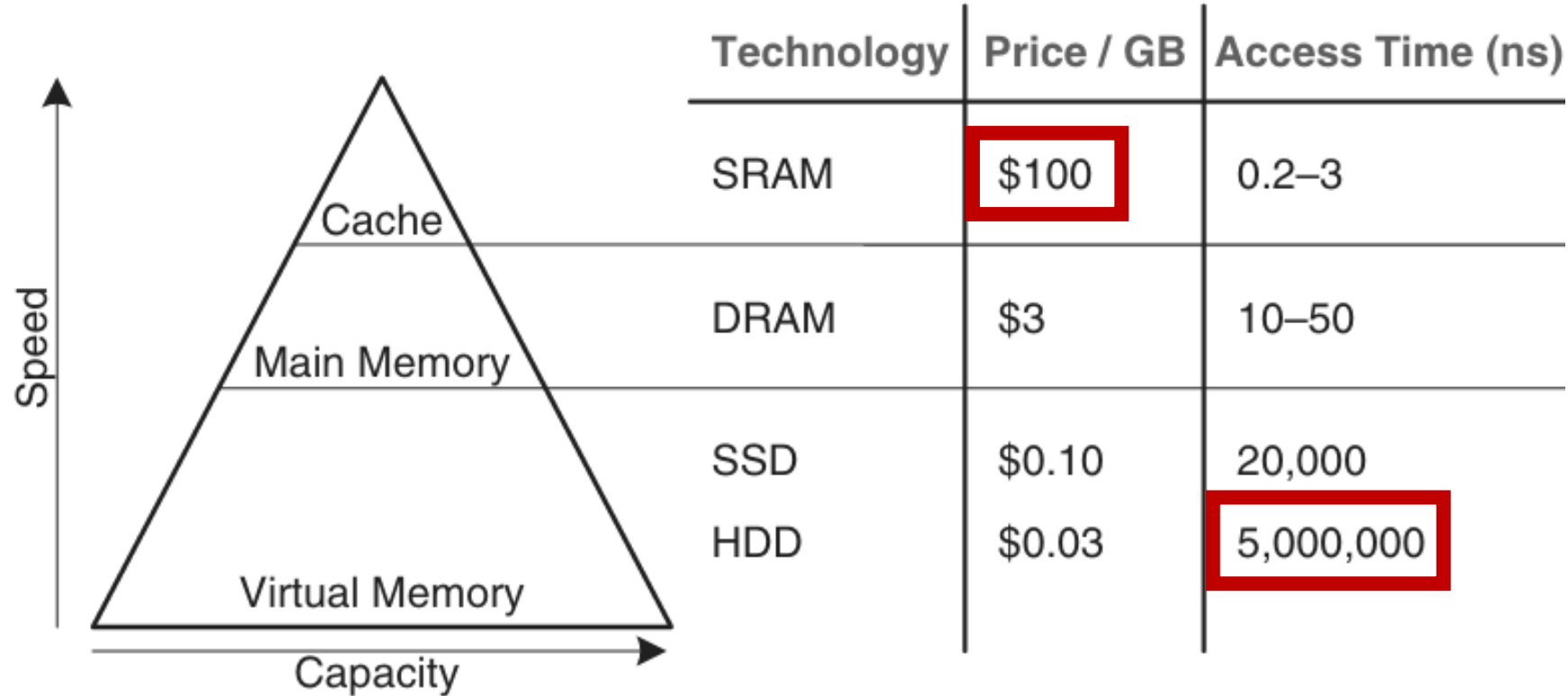
Integer		FP		Memory	
Add		FAdd		Cache	(64bit)
8 bit	0.03pJ	16 bit	0.4pJ	8KB	10pJ
32 bit	0.1pJ	32 bit	0.9pJ	32KB	20pJ
Mult		FMult		1MB	100pJ
8 bit	0.2pJ	16 bit	1.1pJ	DRAM	1.3-2.6nJ
32 bit	3.1pJ	32 bit	3.7pJ		



10,000x

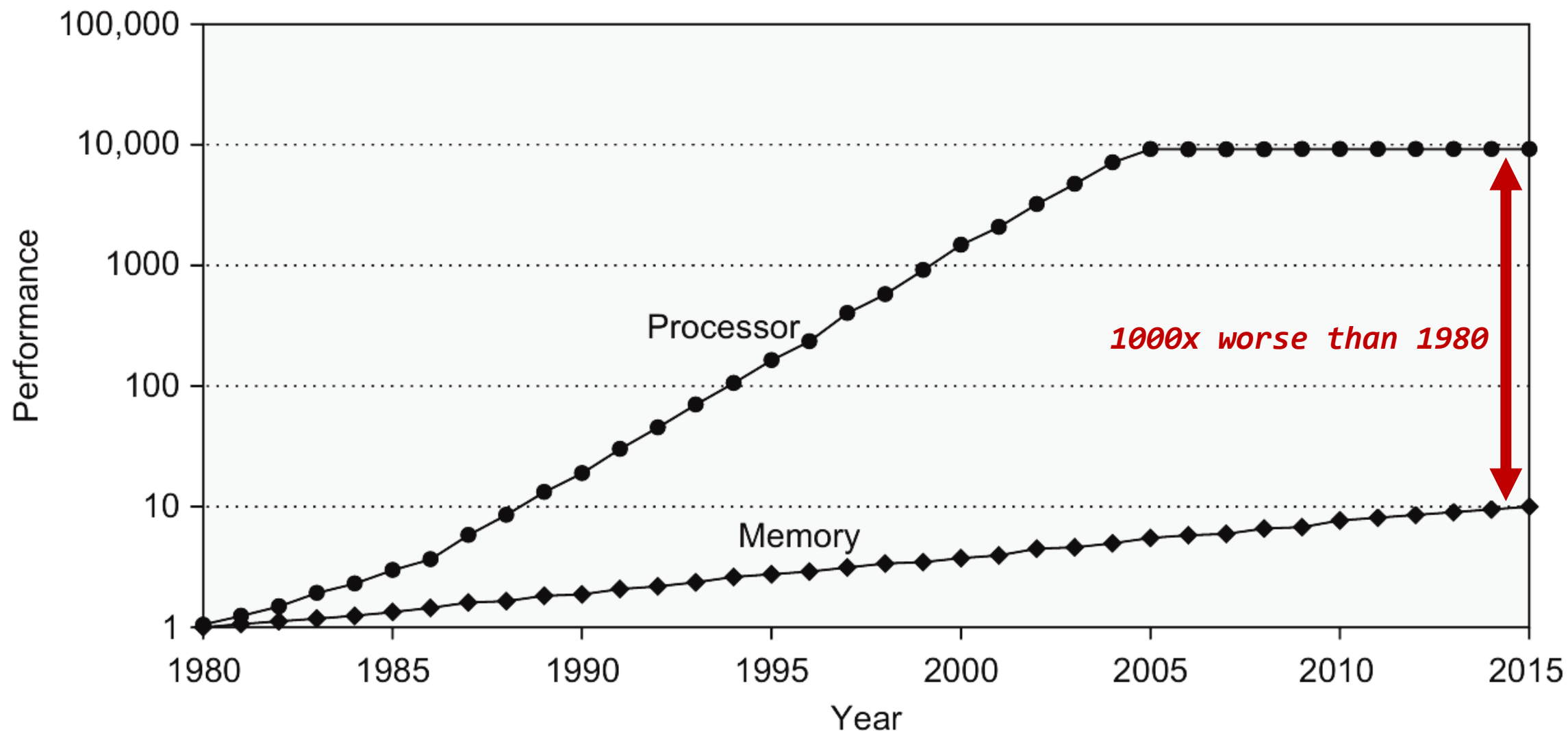
Horowitz, Mark. "Computing's Energy Problem (and What We Can Do About It)," ISSCC, 2014.

Cost-Performance Tradeoff



- There are **other concerns** too (e.g., energy, reliability, etc...)

It's Just Hard to Improve Memory



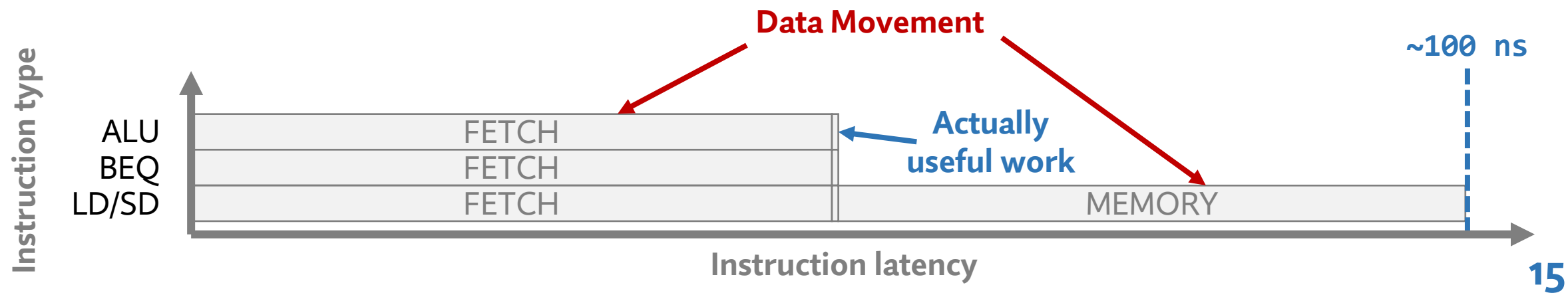
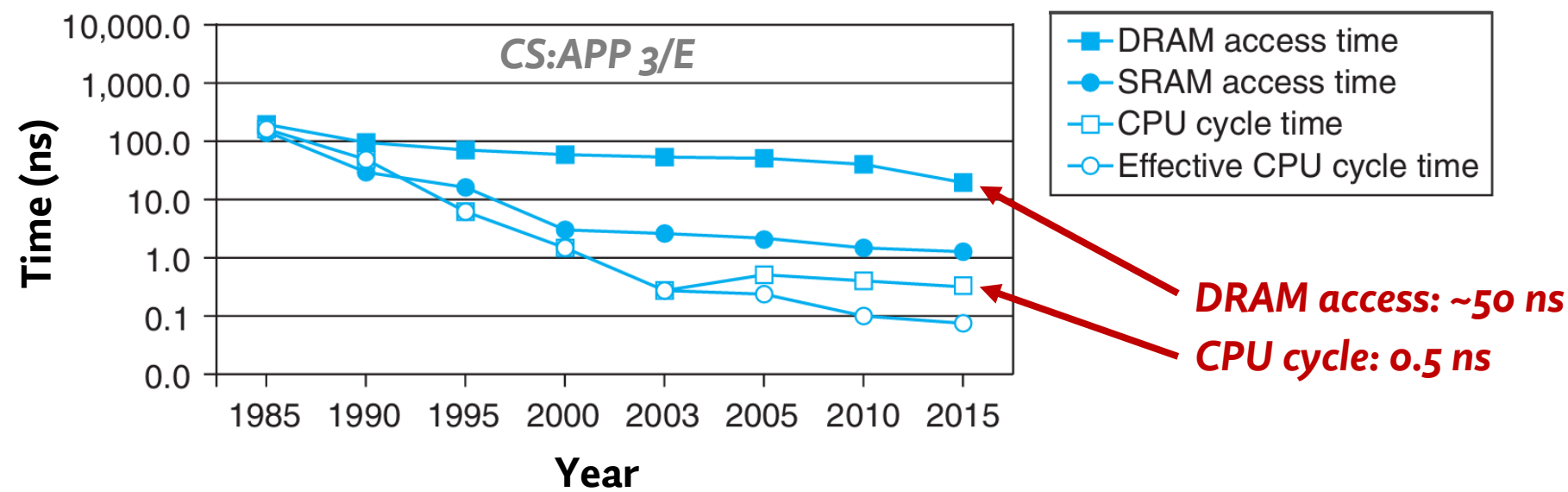
Hennessy & Patterson, "A Quantitative Approach to Computer Architecture," 6/E., Figure 2.2

Agenda

- Why Is Main Memory Slow?
- **Caching and the Memory Hierarchy**
- Direct-Mapped Caches
 - Where to Place Data
 - How to Query the Cache

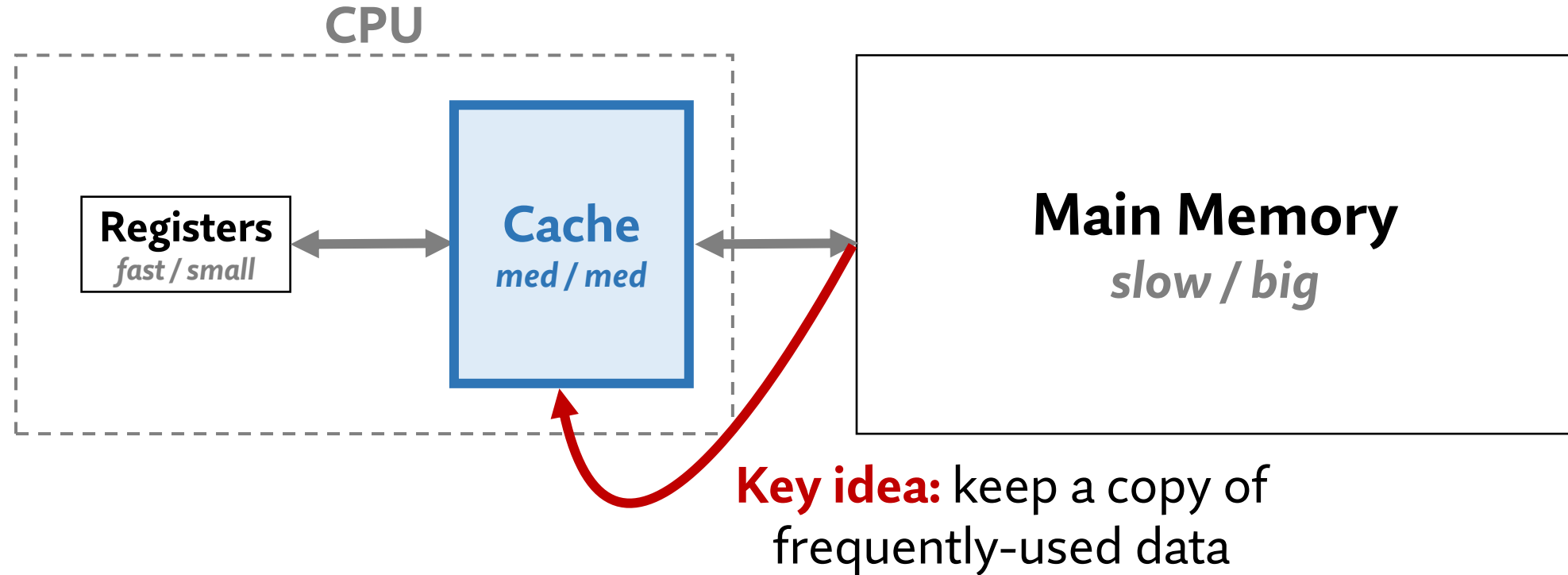
Main Memory Kills Performance

- The CPU wastes lots of time **just moving data**



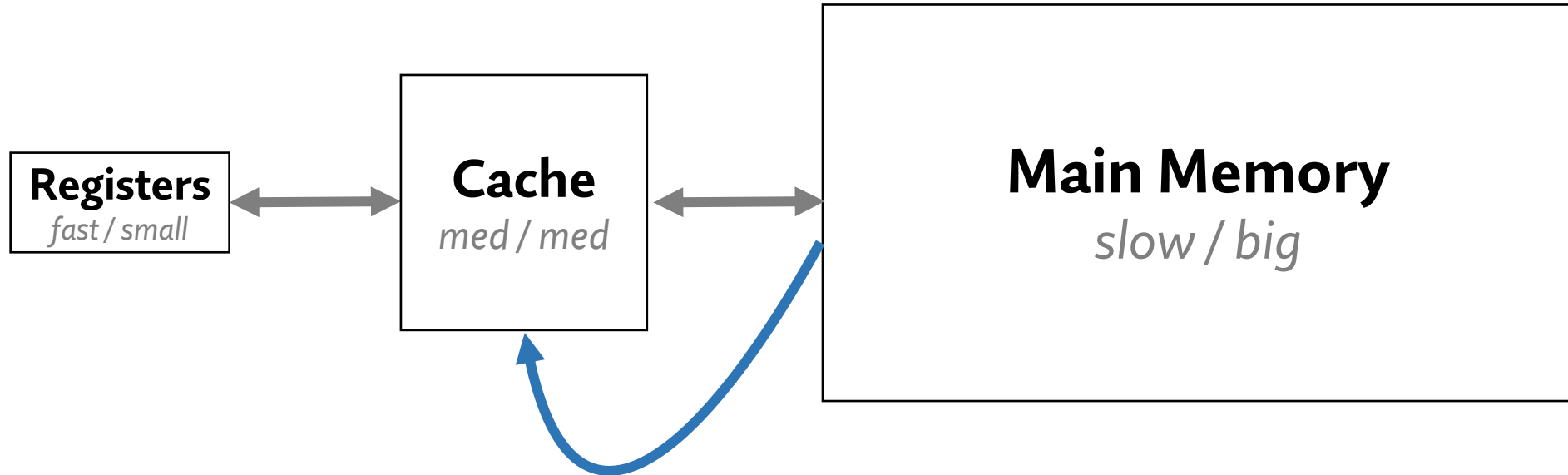
Getting the Best of Both Worlds

- **Goal:** Give the **illusion** that main memory is **large but fast**



- If we cache **the right data**, the CPU never goes to main memory
 - **Common case:** accesses are at the speed of cache
 - **Worst case:** data not in the cache

The Right Data to Cache



Million-dollar question:
What should we cache?

- **A:** Data that we know we will need soon (prediction problem!)
 - **Cache hit:** data is in the cache - fast
 - **Cache miss:** data is NOT in the cache - slow

Choosing What to Cache

- The answer comes from studying application behavior

```
void func(uint64_t *a, uint64_t N)
{
    for(int i = 0; i < N; i++)
        a[i]++;
}
```

```
func:
    slli a1, a1, 3
    add a1, a0, a1
.L2:
    bne a0, a1, .L3
    ret
.L3:
    ld a5, 0(a0)
    addi a0, a0, 8
    addi a5, a5, 1
    sd a5, -8(a0)
    j .L2
```

Memory Access Pattern

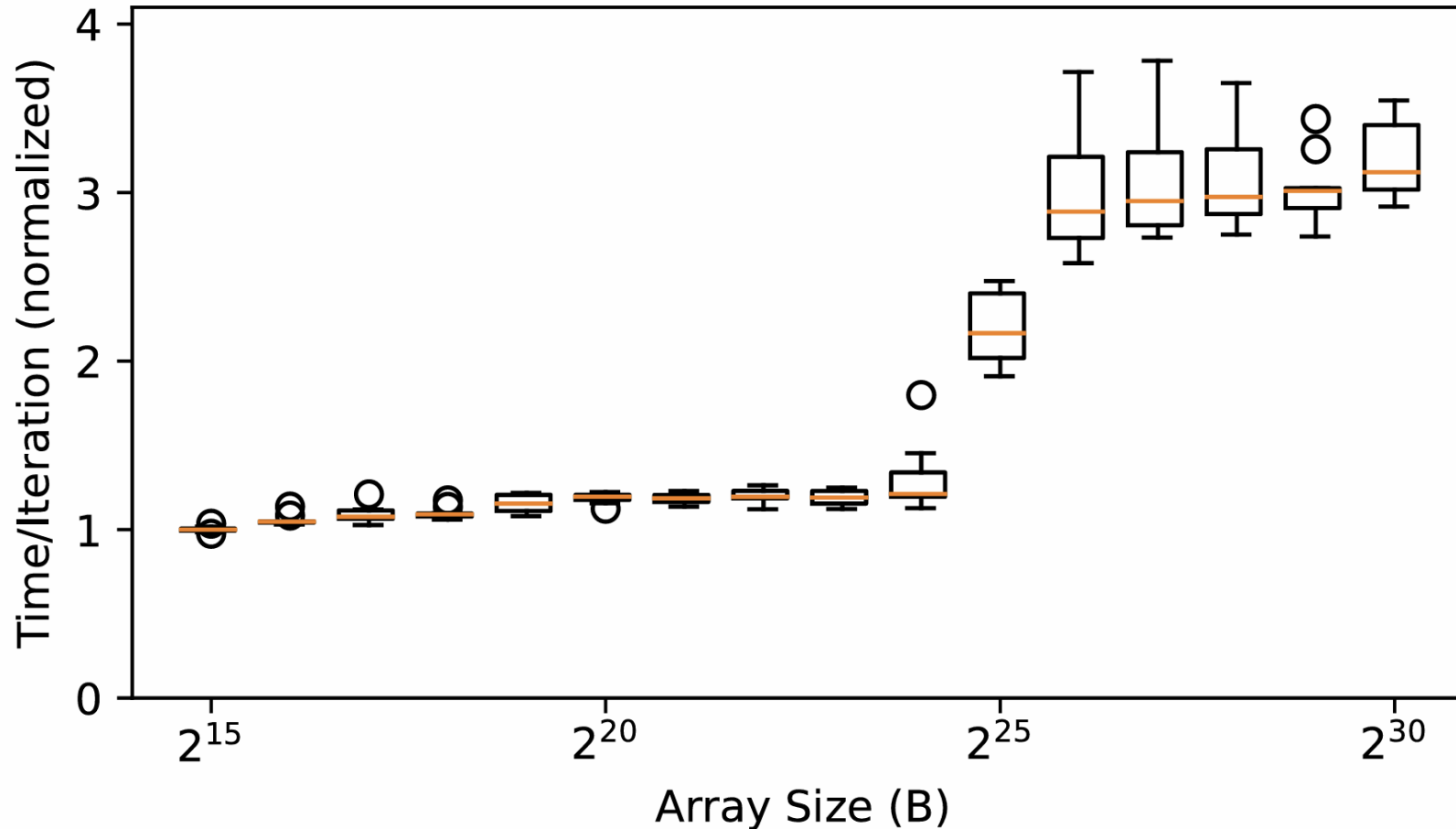
```
read(A[0])
write(A[0])
read(A[1])
write(A[1])
...
read(A[N-1])
write(A[N-1])
```

Observation 1: Re-use

Observation 2: Adjacency

Benchmarking ilab2

```
for(uint64_t i = 0; i < array_size; i += stride)
    temp = arr[i];
```



```
li    a4, 0
.loop:
    bge    a4, a2, .done
    ld     t1, 0(t0)
    add    a4, a4, a3
    add    t0, t0, a3
    j      .loop
.done:
    ...
```

Example: Matrix Multiply ($C = A*B$)

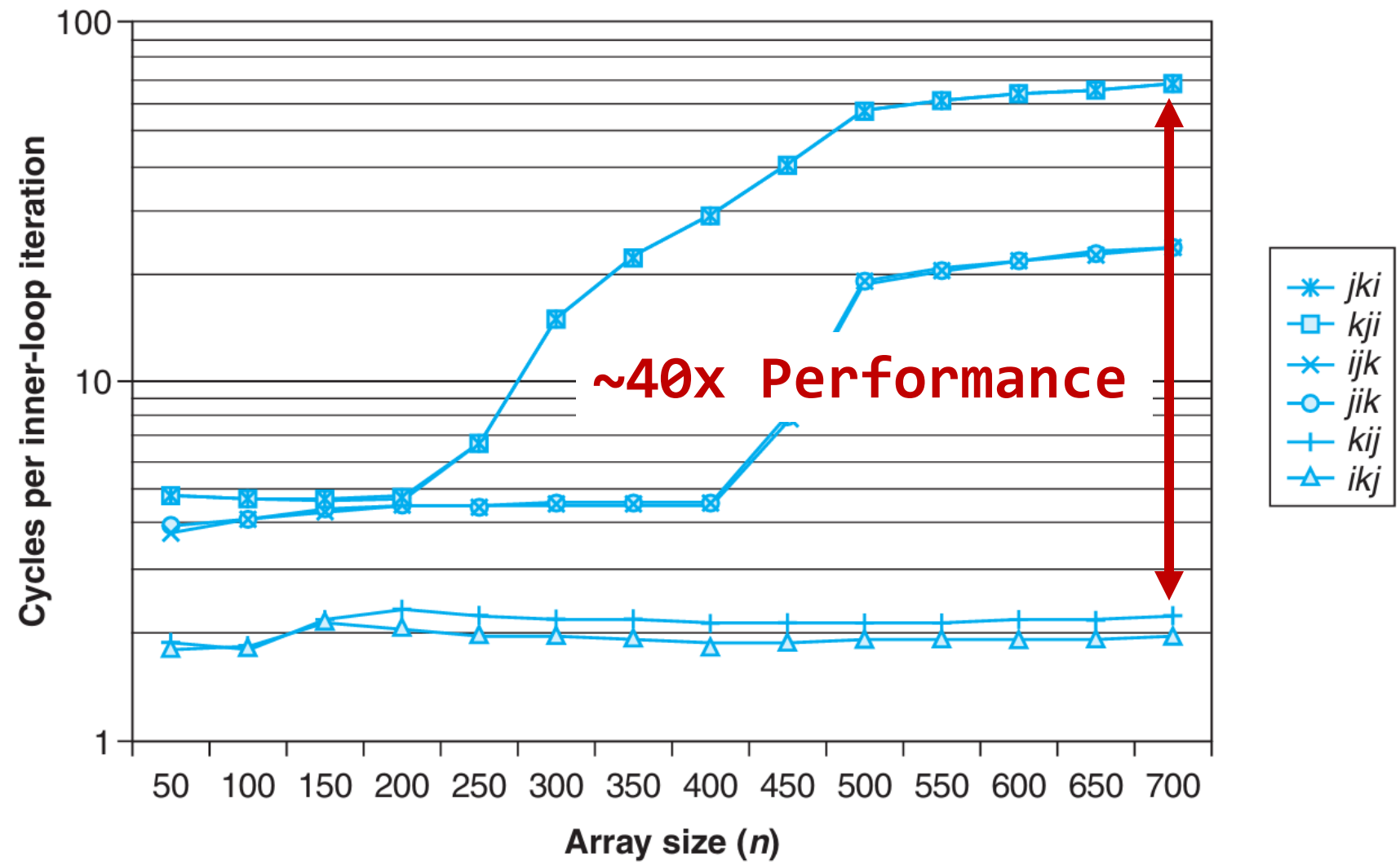


Figure 6.46 Core i7 matrix multiply performance.
CS:APP 3/E, “Section 6.6.2”

(a) Version *ijk* code/mem/matmult/mm.c

```
1 for (i = 0; i < n; i++)
2     for (j = 0; j < n; j++) {
3         sum = 0.0;
4         for (k = 0; k < n; k++)
5             sum += A[i][k]*B[k][j];
6         C[i][j] += sum;
7     }
```

code/mem/matmult/mm.c

(c) Version *jki* code/mem/matmult/mm.c

```
1 for (j = 0; j < n; j++)
2     for (k = 0; k < n; k++) {
3         r = B[k][j];
4         for (i = 0; i < n; i++)
5             C[i][j] += A[i][k]*r;
6     }
```

code/mem/matmult/mm.c

(e) Version *kij* code/mem/matmult/mm.c

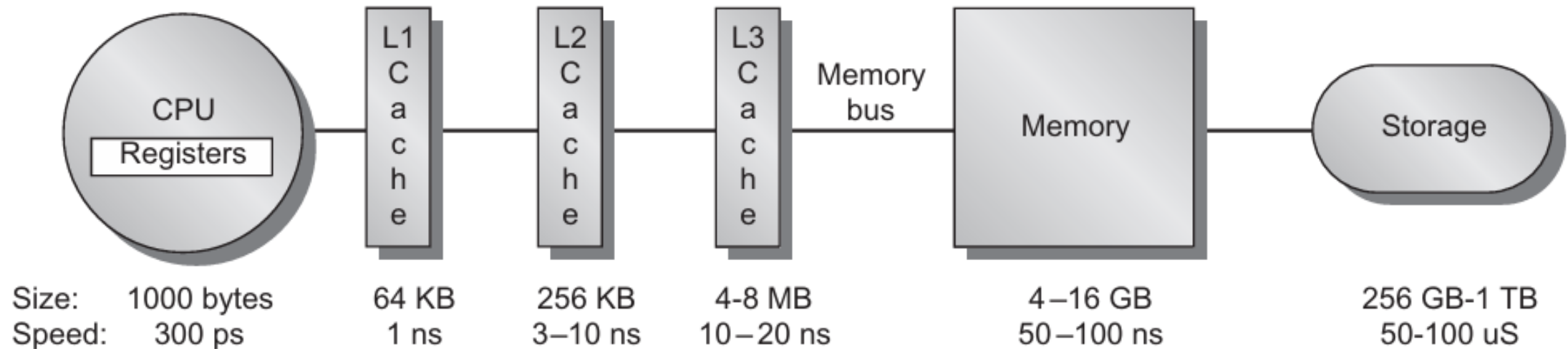
```
1 for (k = 0; k < n; k++)
2     for (i = 0; i < n; i++) {
3         r = A[i][k];
4         for (j = 0; j < n; j++)
5             C[i][j] += r*B[k][j];
6     }
```

code/mem/matmult/mm.c

Figure 6.44 Six versions of matrix multiply. Each
20

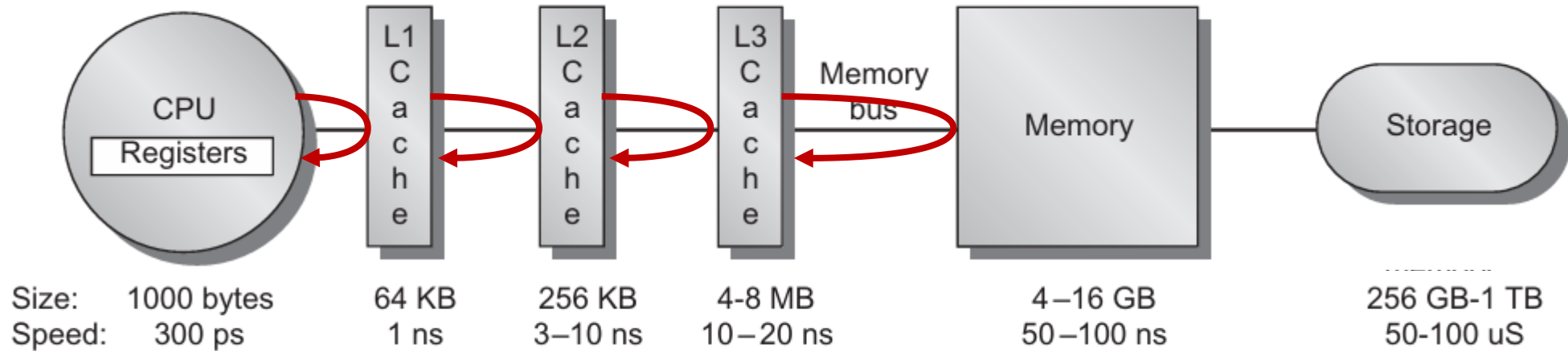
The Memory Hierarchy

- **Key idea:** Have successively larger caches
 - “Hot data” close to CPU (closer the better)
 - “Cold data” far from CPU



Hit/Miss Performance

- Performance depends on whether there is a **cache hit or miss**



- L1 hit:
- L1 miss + L2 hit:
- L1 miss + L2 miss + L3 hit:
- L1 miss + L2 miss + L3 miss + Memory:

Testing Memory/Cache Latency

AIDA64 Cache & Memory Benchmark

	Read <i>i</i>	Write <i>i</i>	Copy <i>i</i>	Latency
Memory	103.47 GB/s	86524 MB/s	96072 MB/s	62.1 ns
L1 Cache	6257.3 GB/s	3141.2 GB/s	6261.3 GB/s	0.9 ns
L2 Cache	1783.2 GB/s	881.14 GB/s	1365.2 GB/s	3.2 ns
L3 Cache	183.67 GB/s	133.86 GB/s	161.95 GB/s	19.3 ns
CPU Type	12-Core Intel Core i9-7920X (Skylake-X, LGA2066)			
CPU Stepping	M0			
CPU Clock	4499.7 MHz (original: 2900 MHz, overclock: 55%)			
CPU FSB	100.0 MHz (original: 100 MHz)			
CPU Multiplier	45x	North Bridge Clock		2999.8 MHz
Memory Bus	1999.8 MHz	DRAM:FSB Ratio		60:3
Memory Type	Quad Channel DDR4-4000 SDRAM (18-18-18-36 CR2)			
Chipset	Intel Kaby Point X299, Intel Skylake-X			
Motherboard	Asus ROG Strix X299-E Gaming II			
BIOS Version	1403			

AIDA64 v6.33.5700 / BenchDLL 4.5.846.8-x64 (c) 1995-2021 FinalWire Ltd.

Save

Start Benchmark

Close

Cache Design Questions

- **Which data** should we cache?
- How do we **access the cache**?
- **Where** in the cache should that data go?
- What happens when the cache **is full**?
- What if we **modify** data in the cache?

Agenda

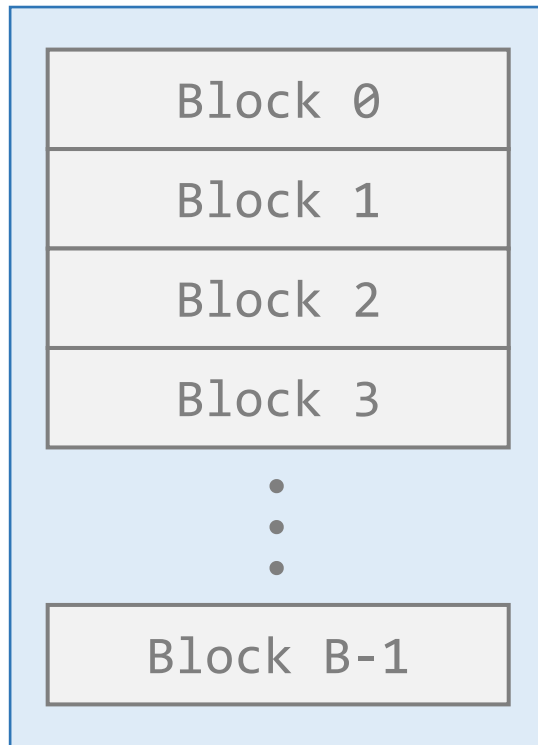
- Why Is Main Memory Slow?
- Caching and the Memory Hierarchy
- **Direct-Mapped Caches**
 - Where to Place Data
 - How to Query the Cache

Cache Blocks

- Caches store “**blocks**” of data (usually 64 bytes)
- Think about it like a smaller version of main memory

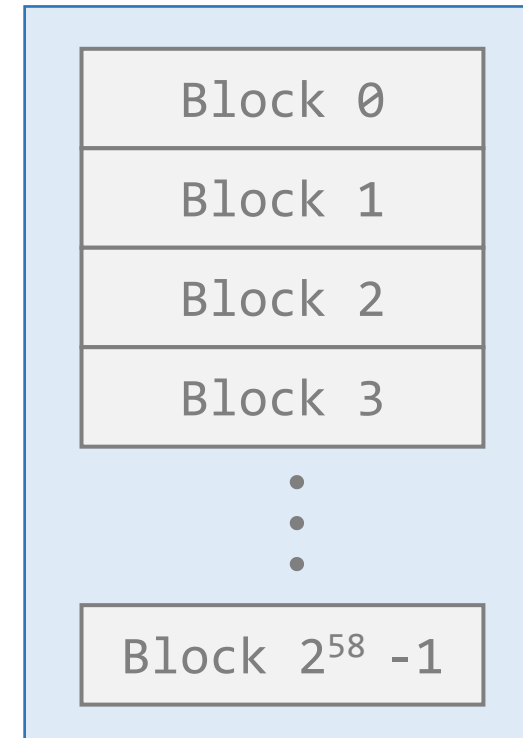
Cache

B blocks



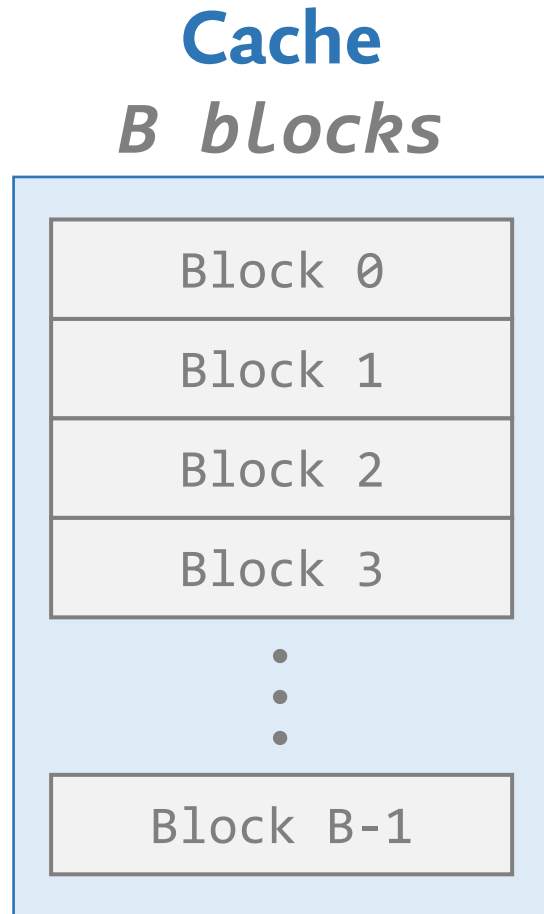
Main Memory

2^{64} / 64 blocks



Types of Caches

- **Q:** Suppose we issue an **ld** (8B): **which block** does it go in?



Fully Associative: Anywhere

Set-Associative: Specific places

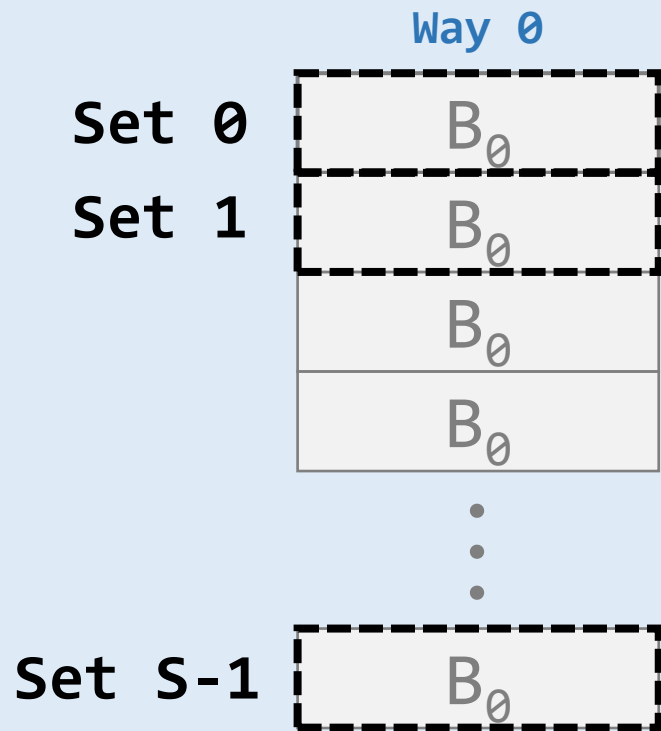
Direct Mapped: Exactly one place

Today

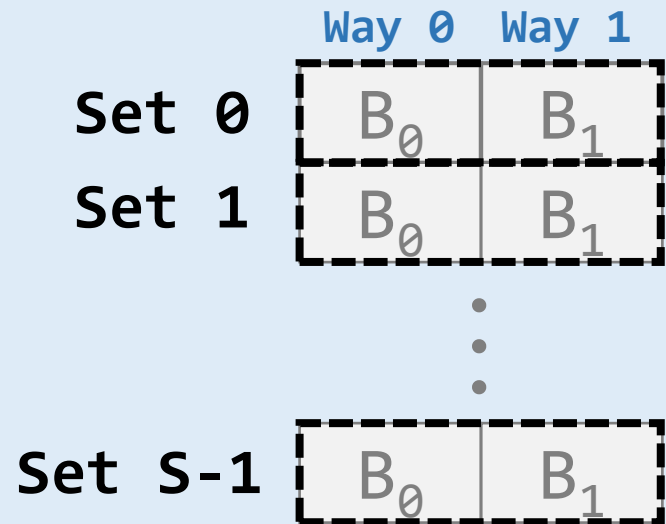
Cache Organizations: Terminology

- Rows (**sets**) and columns (**ways**) of blocks

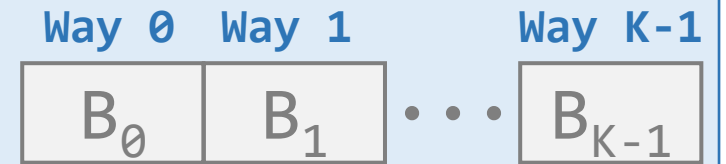
Direct-Mapped



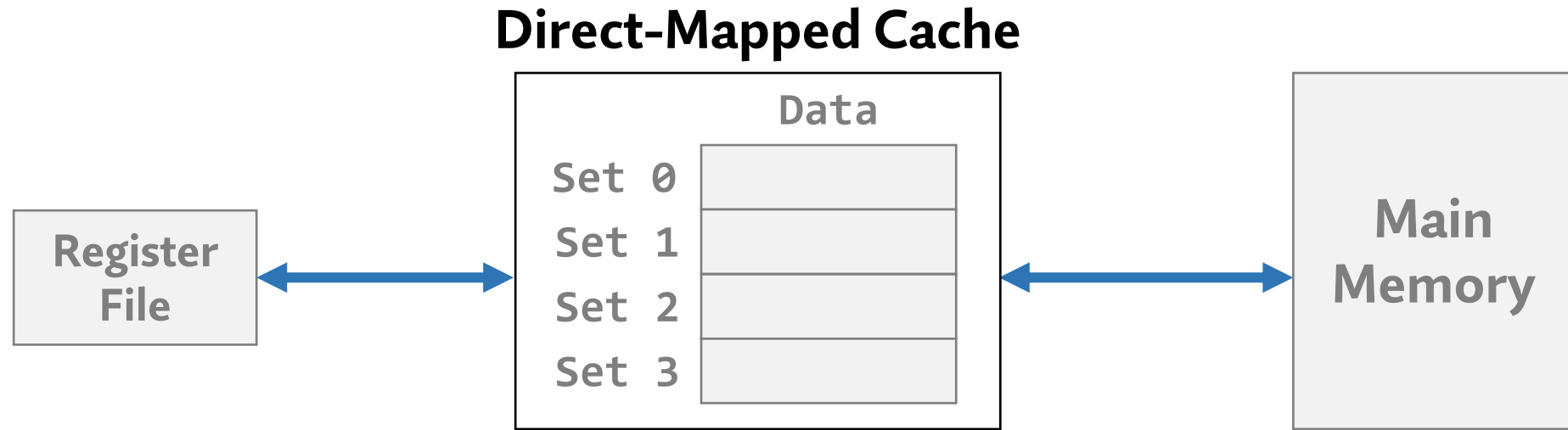
Set-Associative



Fully-Associative



Toward a Direct-Mapped Cache



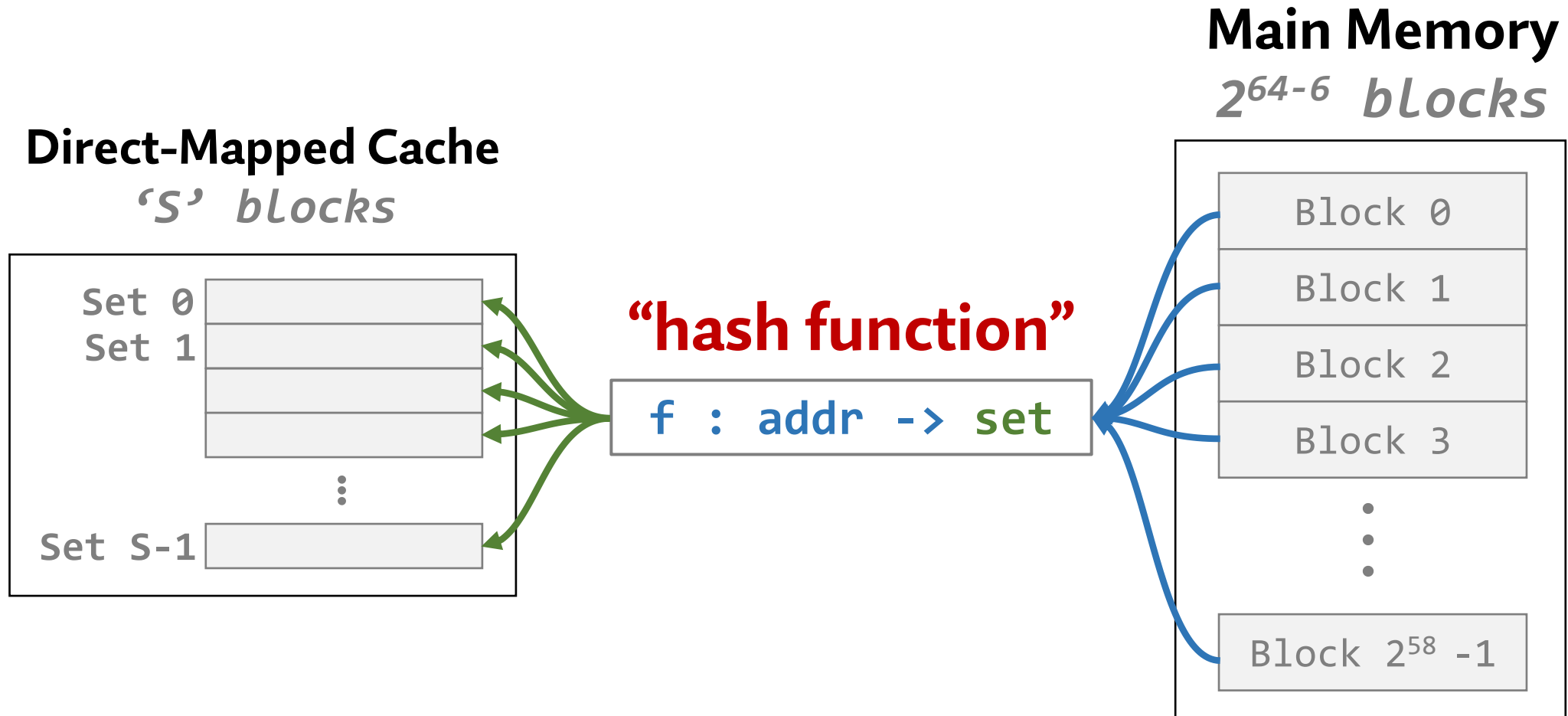
- **Design questions:**
 - Where to place data?
 - How to query the cache?

Agenda

- Why Is Main Memory Slow?
- Caching and the Memory Hierarchy
- Direct-Mapped Caches
 - **Where to Place Data**
 - How to Query the Cache

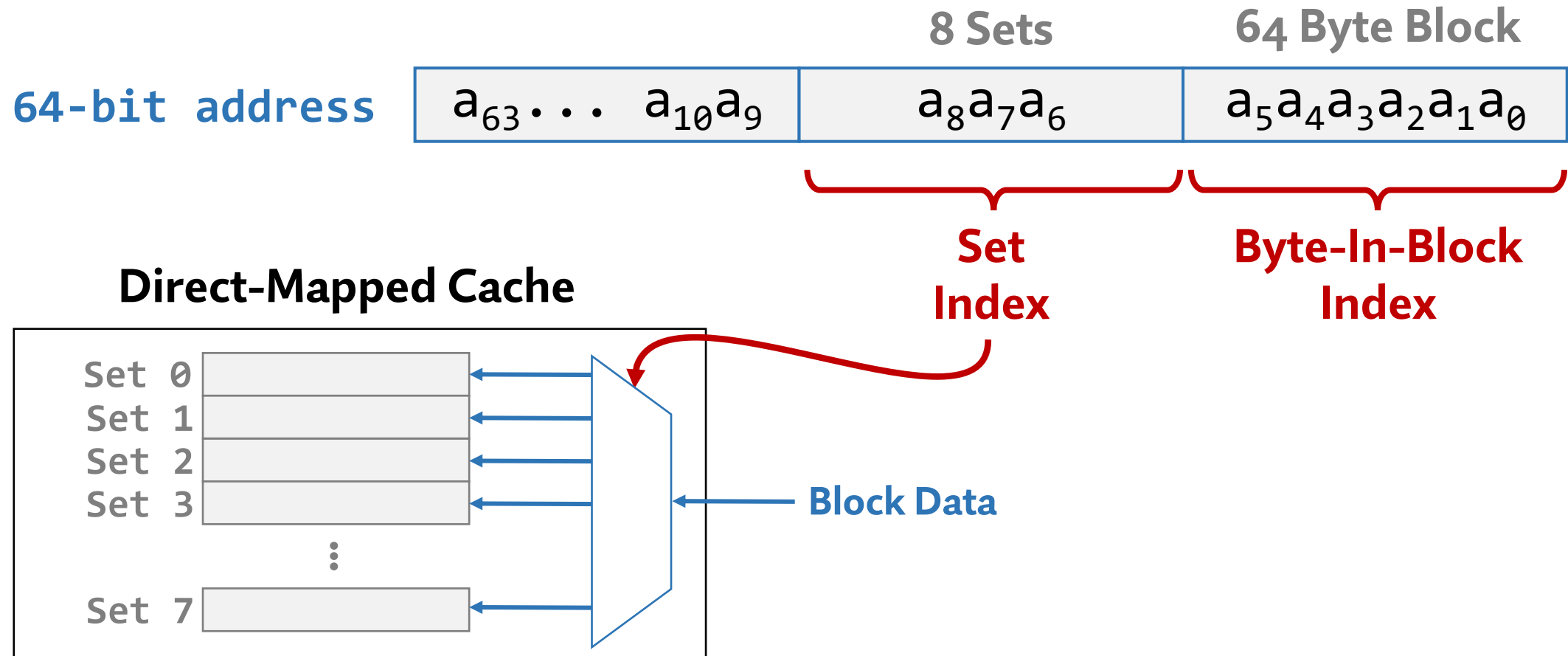
Mapping Data to a Direct-Mapped Cache

- Need to map **every memory location** to a specific set



Set Index Hash Function

- We could use something fancy (lots of research on this!)
- **Simplest:** just use bits of the address



Visualizing the Cache Mapping

Direct-Mapped Cache 4 blocks

Set 0	Block 0
Set 1	Block 1
Set 2	Block 2
Set 3	Block 3

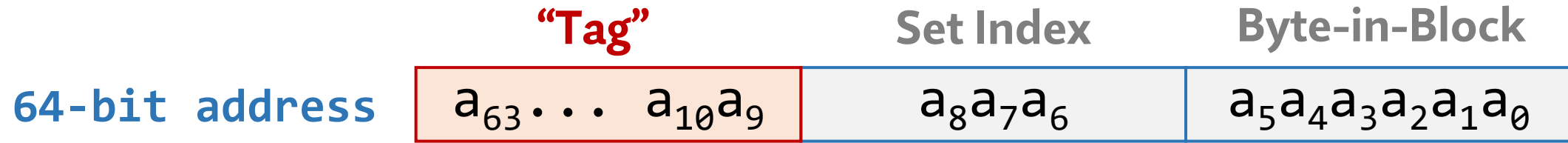
Main Memory 2^{64-6} blocks

Block 0
Block 1
Block 2
Block 3
Block 4
Block 5
Block 6
Block 7
⋮
Block $2^{58} - 1$

**But hash
collision!**

Resolving Hash Collisions

- **Solution:** cache **both data and address**



Direct-Mapped Cache

	Data	Tag
Set 0		
Set 1		
Set 2		
Set 3		
	⋮	
Set 7		

Must **compare tag bits**
on every cache access
to check for collision

Agenda

- Why Is Main Memory Slow?
- Caching and the Memory Hierarchy
- Direct-Mapped Caches
 - Where to Place Data
 - **How to Query the Cache**

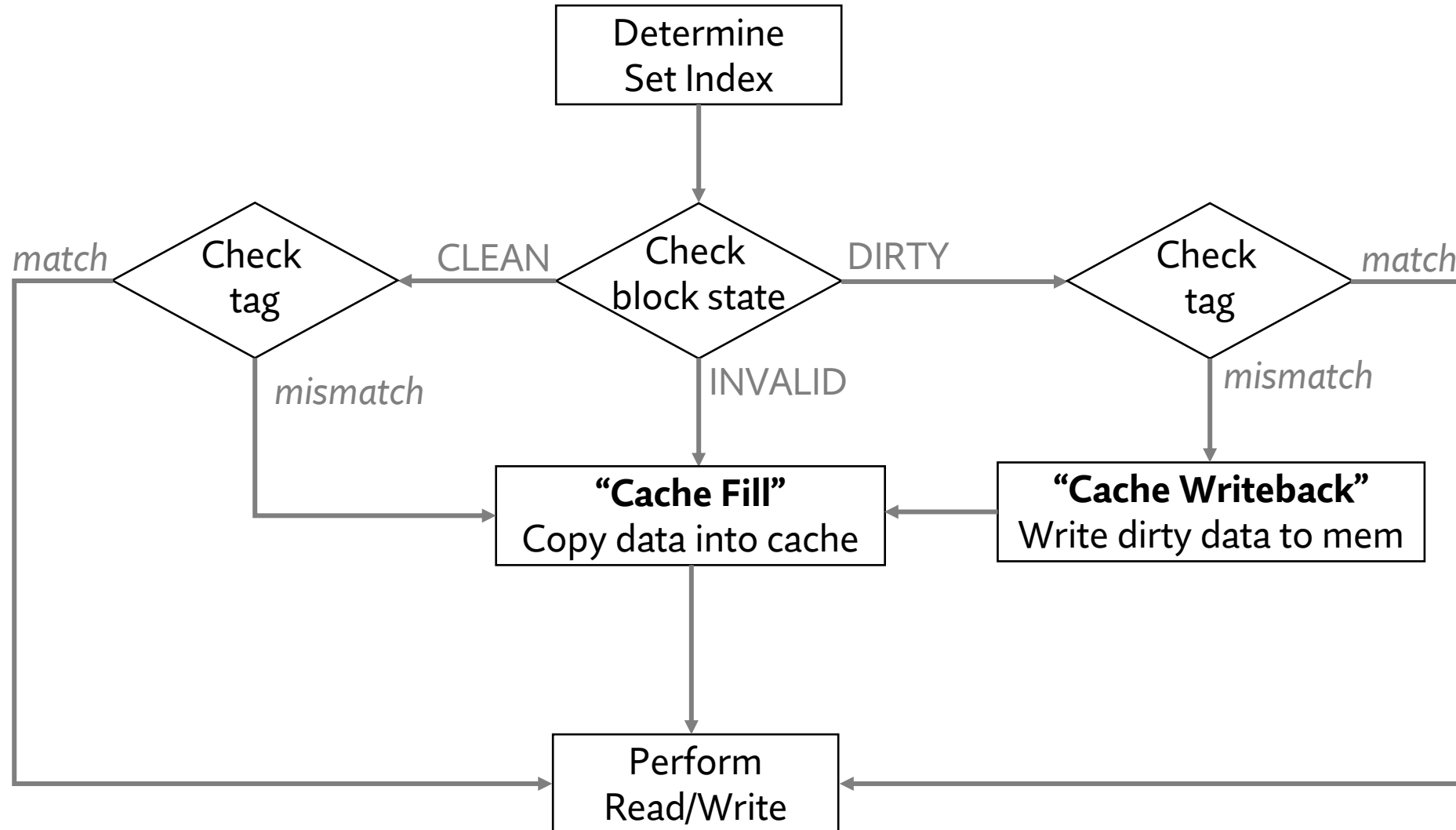
Querying the Cache

- The cache must **track the current state** of every block
 - **INVALID**: no data currently cached (tag is meaningless)
 - **CLEAN**: data is cached and unmodified
 - **DIRTY**: data is cached and modified by the CPU (doesn't match memory)

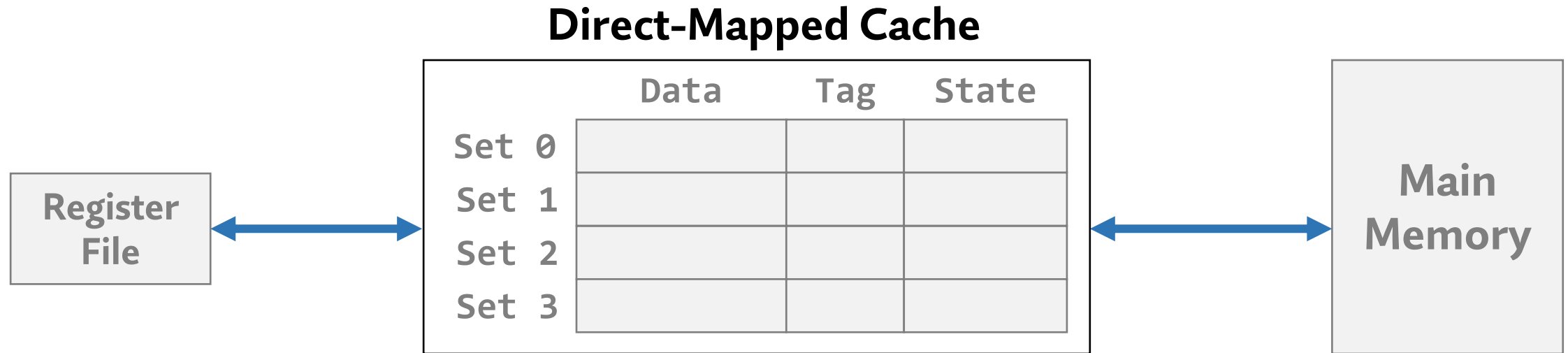
Direct-Mapped Cache

	Data	Tag	State
Set 0			
Set 1			
Set 2			
Set 3			

Querying the Cache: Cache Controller



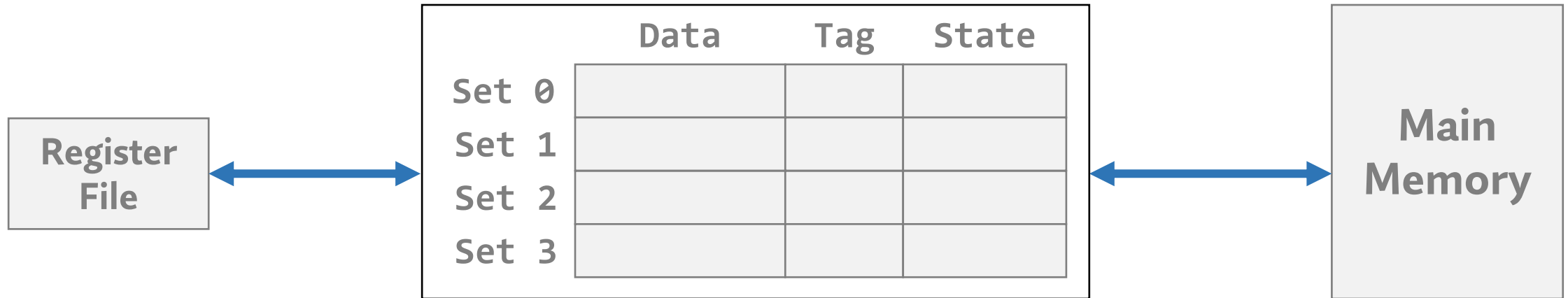
Querying the Cache: STORE Instruction



```
func:
    // a0 = 0x100
    // a1 = 0xff
    sd    a1, 0(a0)
```

Example: Accessing an Array

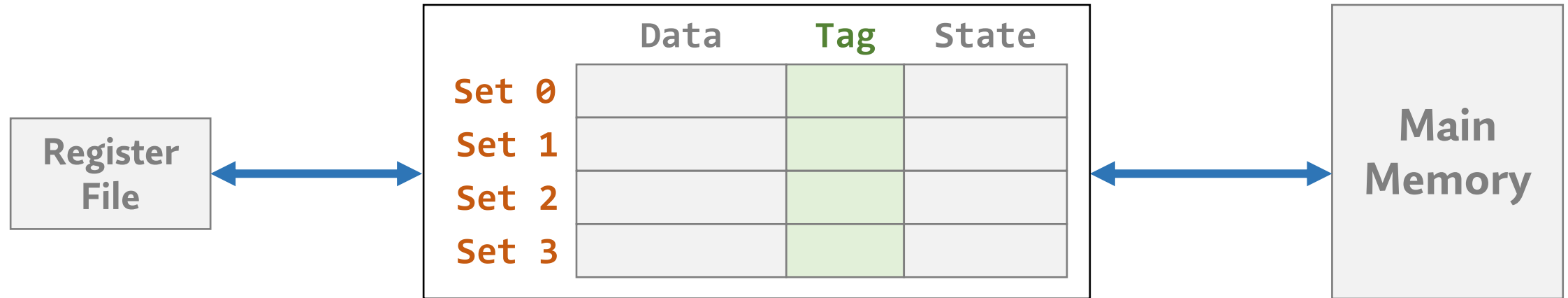
Direct-Mapped Cache



```
void sum(uint64_t a[10])
{
    // a = 0x100
    uint64_t s = 0;
    for(int i = 0; i < 10; i++)
        s += a[i];
    return s;
}
```

Example: Accessing an Array

Direct-Mapped Cache



```
void sum(uint64_t a[10])
{
    // a = 0x100
    uint64_t s = 0;
    for(int i = 0; i < 10; i++)
        s += a[i];
    return s;
}
```

a[0] = 1_0000_0000
a[1] = 1_0000_1000
a[2] = 1_0001_0000
a[3] = 1_0001_1000
a[4] = 1_0010_0000
a[5] = 1_0010_1000
a[6] = 1_0011_0000
a[7] = 1_0011_1000
a[8] = 1_0100_0000
a[9] = 1_0100_1000

read(a[0]) = MISS
read(a[1]) = HIT
read(a[2]) = HIT
read(a[3]) = HIT
read(a[4]) = HIT
read(a[5]) = HIT
read(a[6]) = HIT
read(a[7]) = HIT
read(a[8]) = MISS
read(a[9]) = HIT

CS 211: Intro to Computer Architecture

13.2: Memory Access Patterns and Caching

Minesh Patel

Spring 2025 – Thursday 24 April