

Caches and Memory

CS 3410, Spring 2014

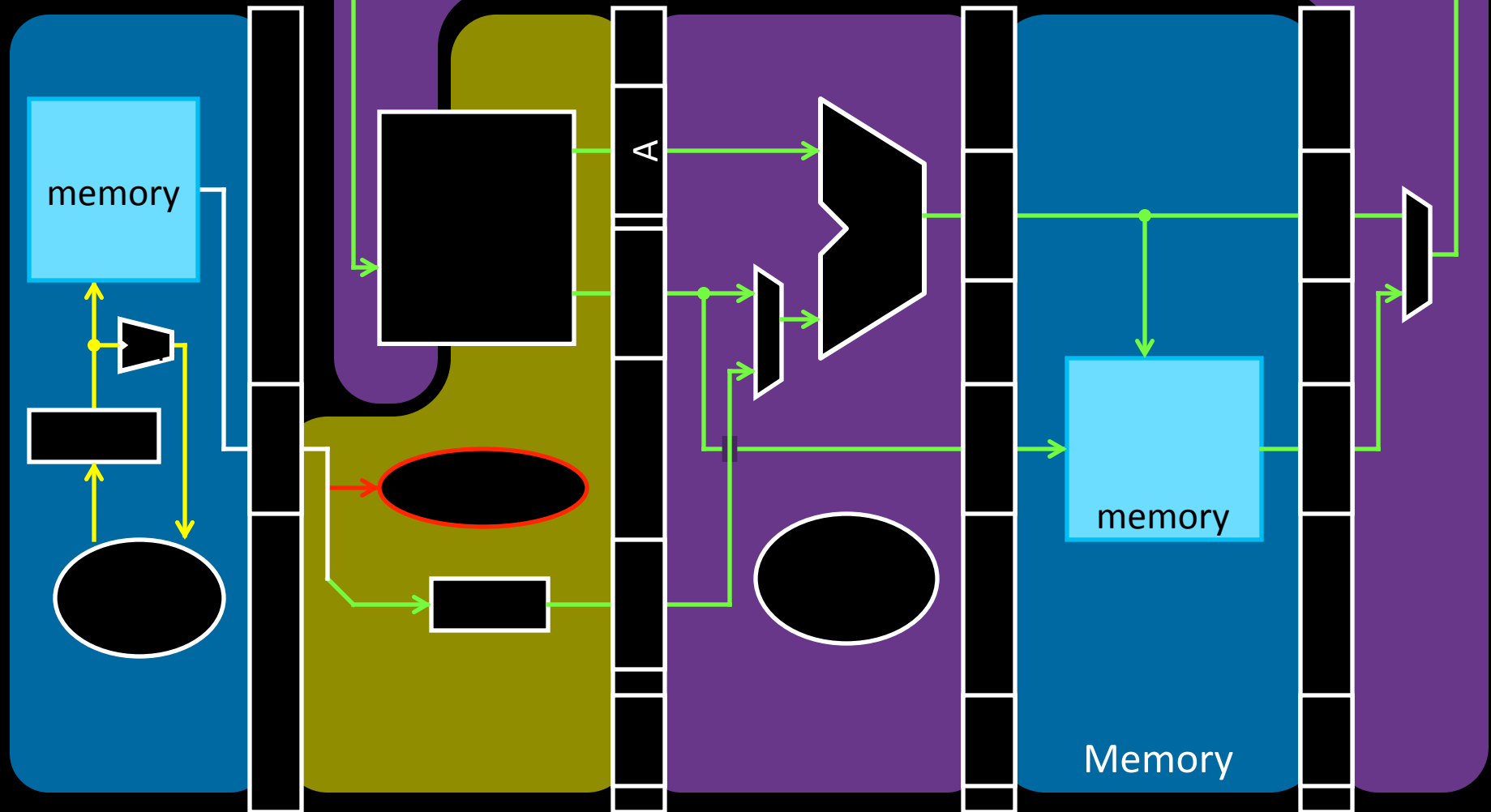
Computer Science

Cornell University

See P&H Chapter: 5.1-5.4, 5.8, 5.15

Pipelined Processor

Code Stored in Memory
(also, data and stack)



Stack, Data, Code
Stored in Memory

Problem

Main memory is very very slow

Remember:

SRAM

- 6-8 transistors, no refresh, fast

DRAM

- 1 transistor, denser, cheaper/bit, needs refresh

Problem

Main memory is very very slow

CPU clock rates $\sim 0.33\text{ns} - 2\text{ns}$ (3GHz-500MHz)

Memory technology	Access time in nanosecs (ns)	Access time in cycles
SRAM (on chip)	0.5-2.5 ns	1-3 cycles
SRAM (off chip)	1.5-30 ns	5-15 cycles
DRAM	50-70 ns	150-200 cycles
SSD (Flash)	5k-50k ns	Tens of thousands
Disk	5M-20M ns	Millions

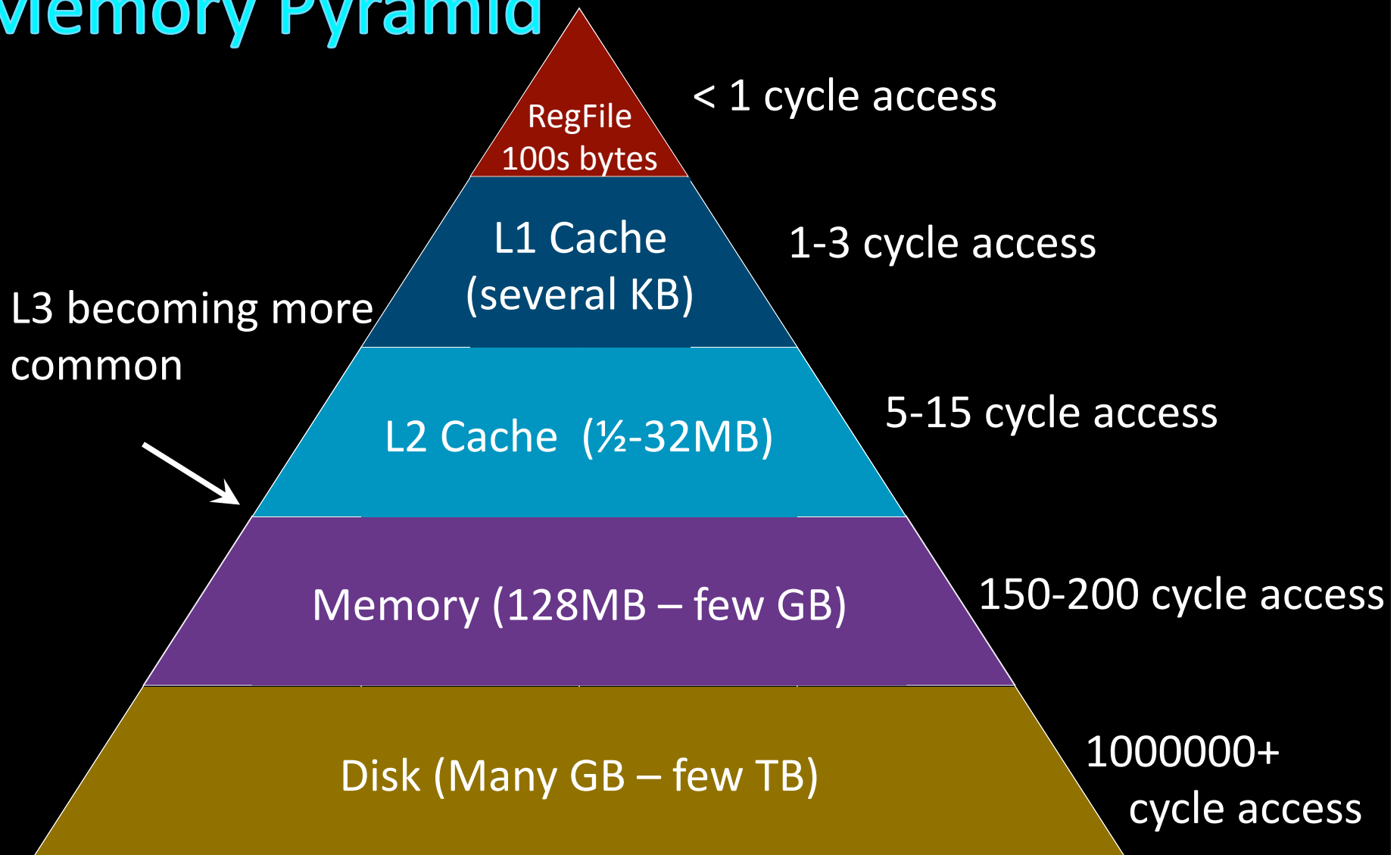
Problem

Main memory is very very slow

CPU clock rates $\sim 0.33\text{ns} - 2\text{ns}$ (3GHz-500MHz)

Memory technology	Access time in nanosecs (ns)	Access time in cycles	\$ per GIB in 2012	Capacity
SRAM (on chip)	0.5-2.5 ns	1-3 cycles		256 KB
SRAM (off chip)	1.5-30 ns	5-15 cycles	\$4k	32 MB
DRAM	50-70 ns	150-200 cycles	\$10-\$20	8 GB
SSD (Flash)	5k-50k ns	Tens of thousands	\$0.75-\$1	512 GB
Disk	5M-20M ns	Millions	\$0.05-\$0.1	4 TB

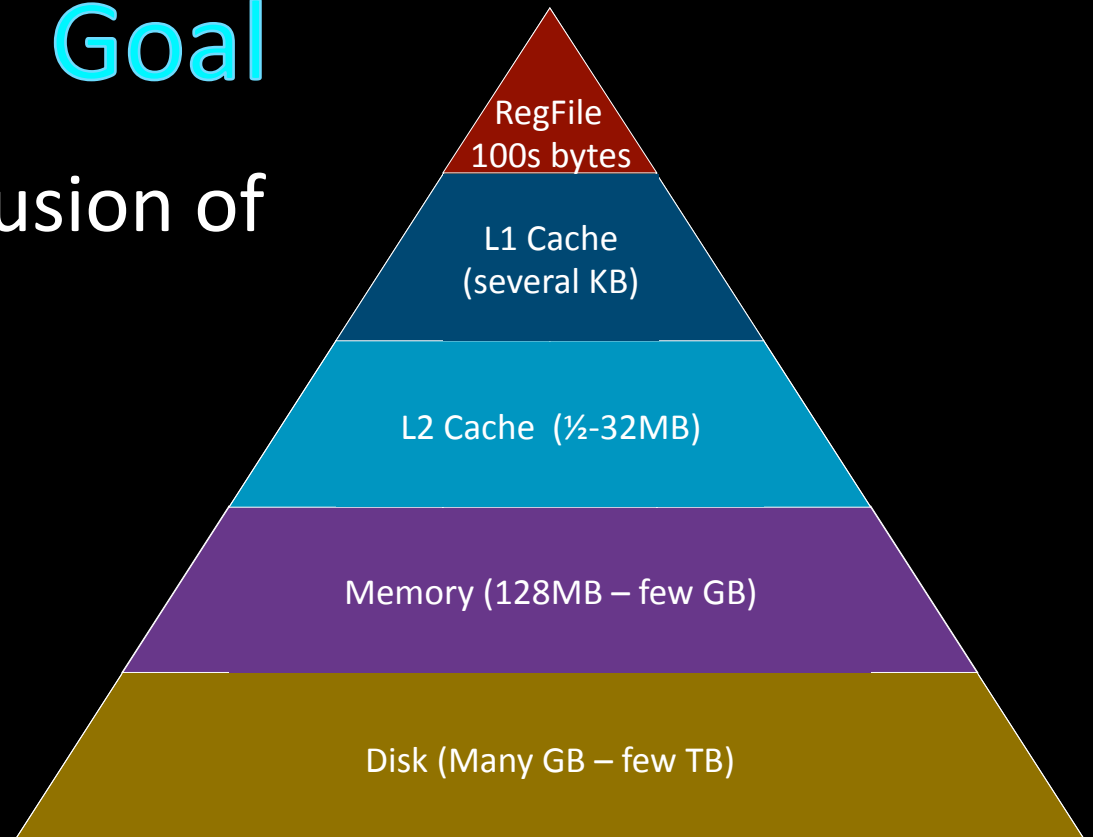
Memory Pyramid



These are rough numbers: mileage may vary for latest/greatest
Caches usually made of SRAM

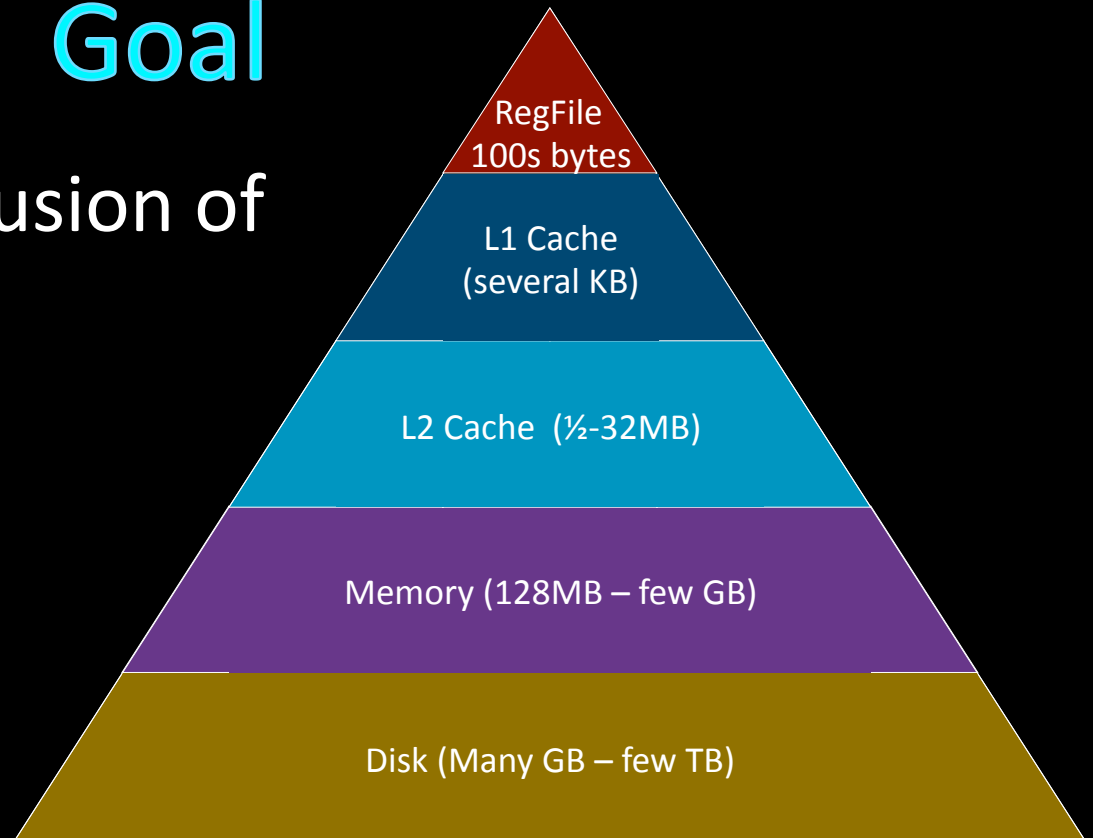
Goal

Can we create an illusion of
cheap,
large and
fast memory?



Goal

Can we create an illusion of
cheap,
large and
fast memory?



Yes, using caches
and assuming temporal and spatial locality

By the end of the cache lectures...

MacBook Pro

Retina, Mid 2012

Processor 2.7 GHz Intel Core i7

Memory 16 GB 1600 MHz DDR3

Graphics NVIDIA GeForce GT 650M 1024 MB

Serial Number C02J70TTDKQ5

Software OS X 10.9.2 (13C64)

Model Name:	MacBook Pro
Model Identifier:	MacBookPro10,1
Processor Name:	Intel Core i7
Processor Speed:	2.7 GHz
Number of Processors:	1
Total Number of Cores:	4
L2 Cache (per Core):	256 KB
L3 Cache:	8 MB
Memory:	16 GB
Boot ROM Version:	MBP101.00EE.B02
SMC Version (system):	2.3f36
Serial Number (system):	C02J70TTDKQ5
Hardware UUID:	F588E08C-60BF-5B35-A087-07714C2B2D11

- 32 KB data + 32 KB instruction **L1 cache** (3 clocks) and 256 KB **L2 cache** (8 clocks) per core.
- Shared L3 cache includes the processor graphics (**LGA 1155**).
- 64-byte **cache** line size.

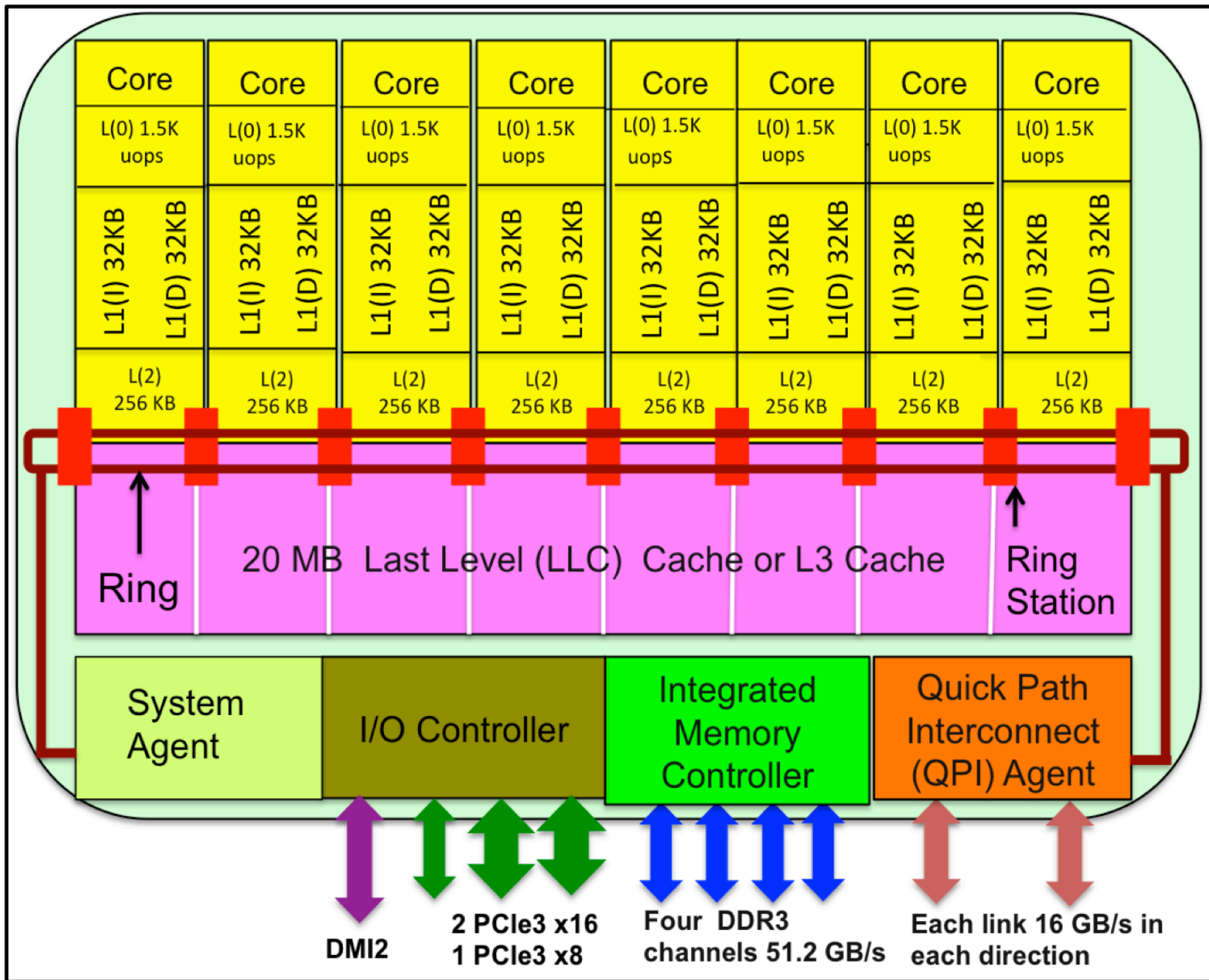
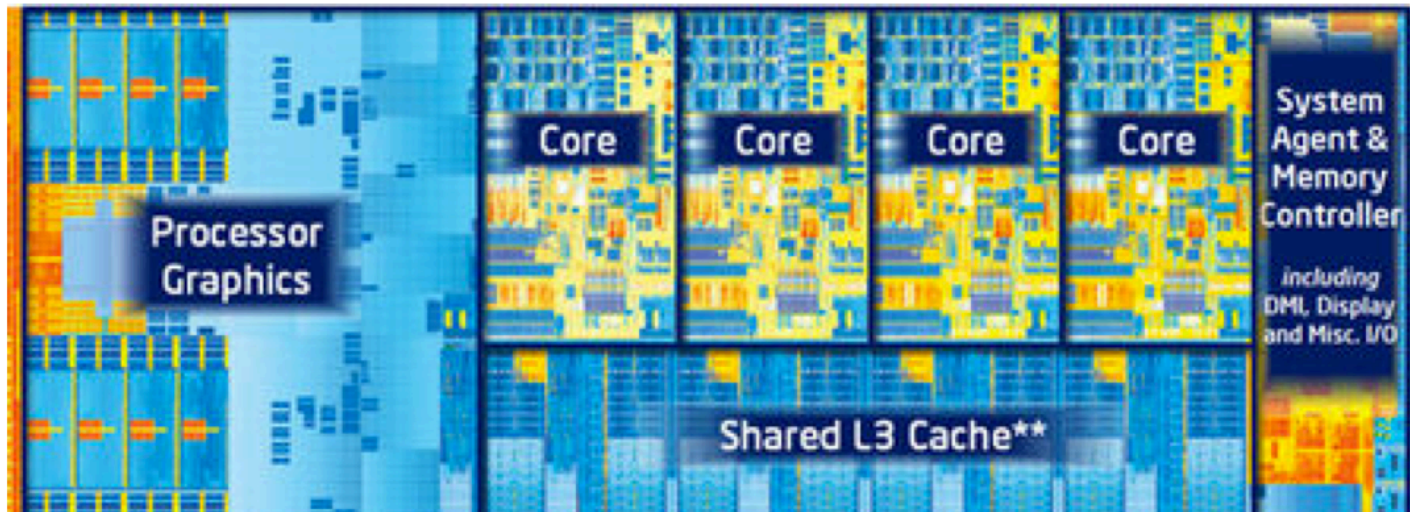
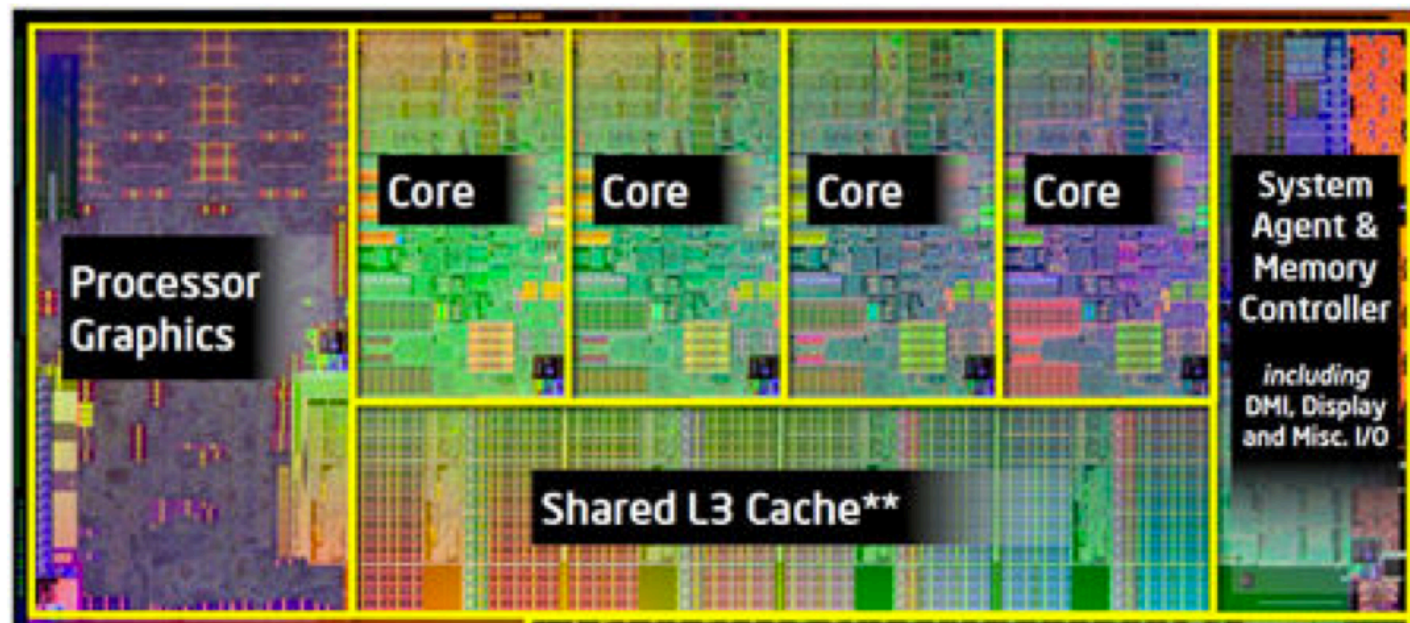


Figure 1. Schematic diagram of a Sandy Bridge processor.



Ivy Bridge: 1.4 billion transistors; 160 square millimeters



Sandy Bridge: 995 million transistors; 216 square millimeters

Today

Caches vs memory vs tertiary storage

- Tradeoffs

Cache organization

- Direct Mapped
- Fully Associative
- N-way set associative

Caching Questions

- How does a cache work? How fast? How big?

Library analogy

Writing a paper on Beren and Lúthien



Book Analogy

- Pick a small set of books; not entire shelf
- Spend time on small set of chapters

Book Analogy

- Pick a small set of books; not entire shelf
- Spend time on small set of chapters
- Sometimes get other books as well
 - Norse mythology, Tolkien biography
- Your desk: out of space
 - Replace less useful books with new ones

Book Analogy

- Pick a small set of books; not entire shelf
 - Cache vs. main memory
 - Working set (the subset in use)
- Spend time on small set of chapters
 - Cache hit
 - Locality of access: temporal and spatial
- Sometimes go to other books
 - Cache may not have data (cache miss)
- Shelf out of space
 - Cache eviction policy

Forms of locality

```
int n = 4;  
int k[] = { 3, 14, 0, 10 };
```

```
int fib(int i) {  
    if (i <= 2) return i;  
    else return fib(i-1)+fib(i-2);  
}
```

Temporal Locality

```
int main(int ac, char **av) {  
    for (int i = 0; i < n; i++) {  
        printi(fib(k[i]));  
        prints("\n");  
    }  
}
```

Spatial Locality

Insight for Caches

If Mem[x] was accessed *recently*...

... then Mem[x] is likely to be accessed *soon*

- Exploit **temporal locality**:
 - Put recently accessed Mem[x] higher in memory hierarchy since it will likely be accessed again soon

... then Mem[x $\pm \epsilon$] is likely to be accessed *soon*

- Exploit **spatial locality**:
 - Put **entire block** containing Mem[x] and surrounding addresses higher in memory hierarchy since nearby address will likely be accessed

Memory Hierarchy

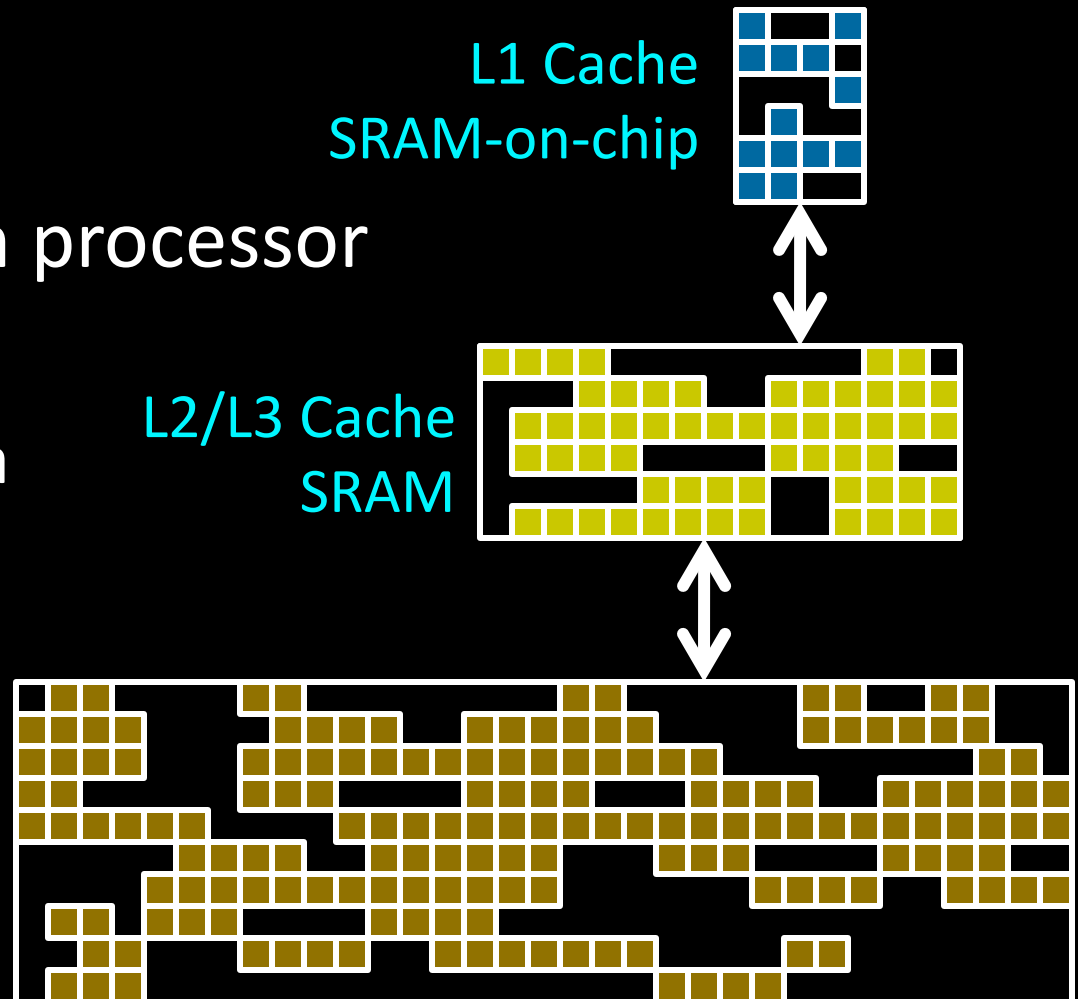
Memory closer to processor

- small & fast
- stores active data

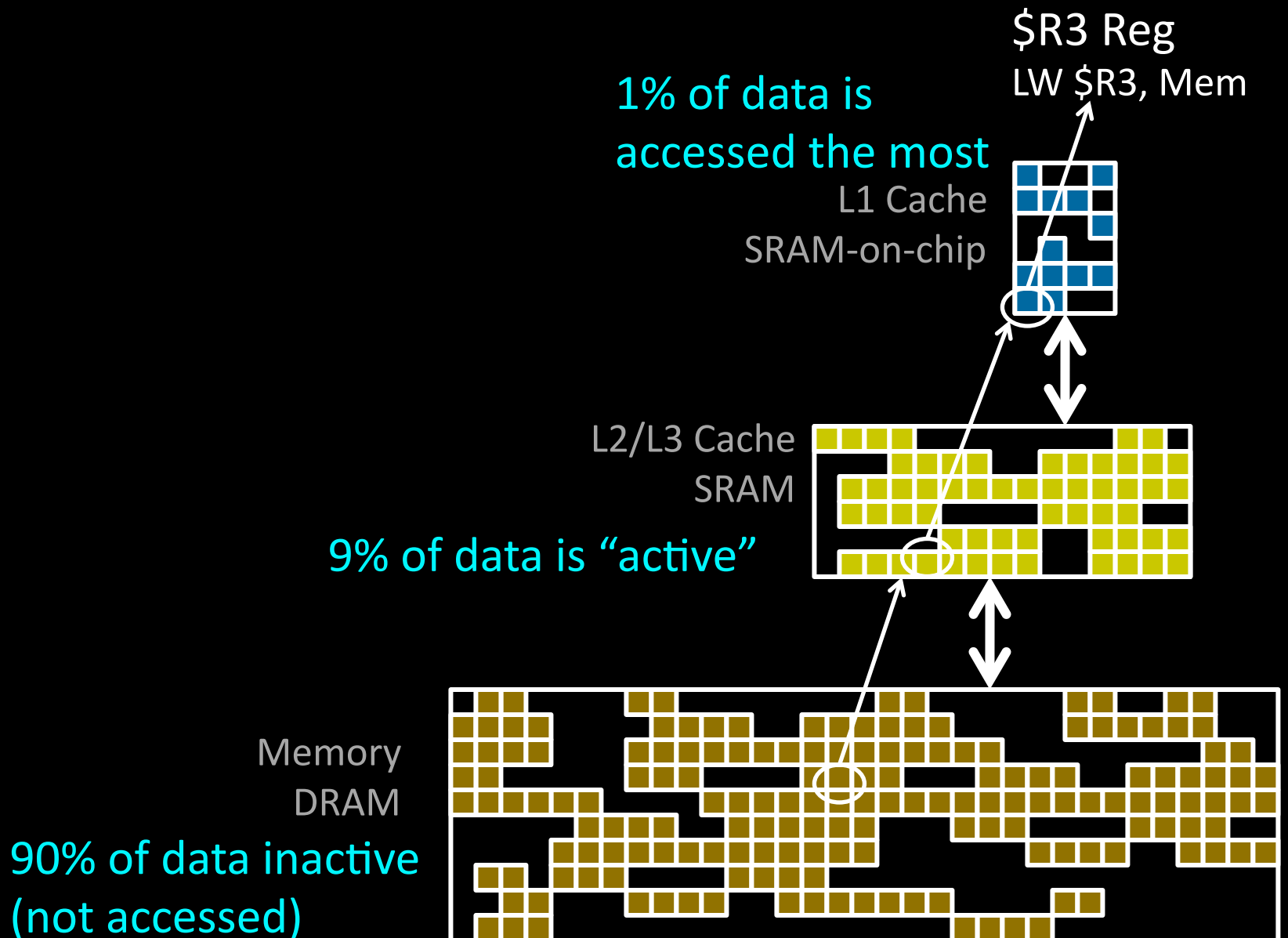
Memory farther from processor

- big & slow
- stores inactive data

Memory
DRAM



90-10 Rule



Memory Hierarchy

Memory closer to processor is fast but small

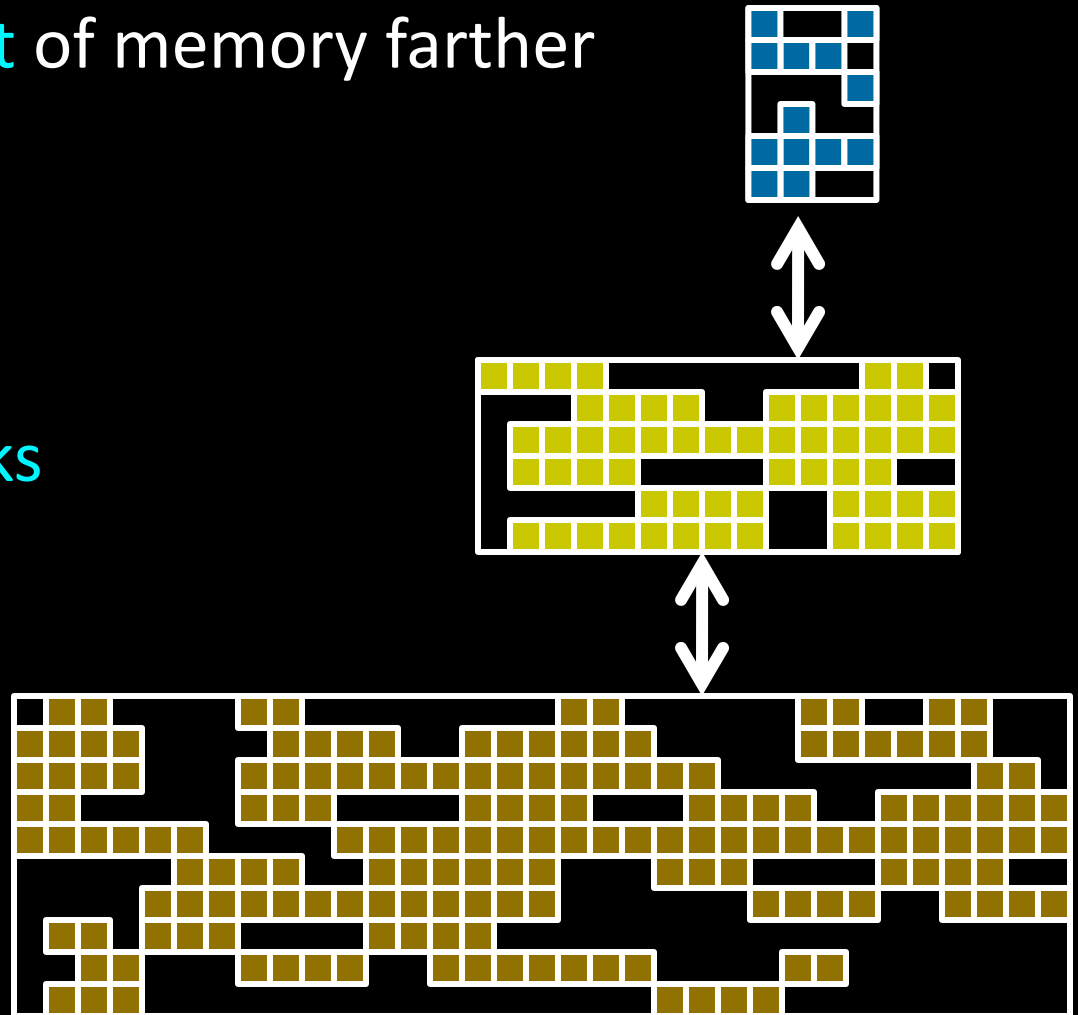
- usually stores **subset** of memory farther
 - “strictly inclusive”

- Transfer whole **blocks** (cache lines):

4kb: disk $\leftrightarrow$ RAM

256b: RAM $\leftrightarrow$ L2

64b: L2 $\leftrightarrow$ L1



Cache Lookups (Read)

Processor tries to access Mem[x]

Check: is block containing Mem[x] in the cache?

- Yes: **cache hit**
 - return requested data from cache line
- No: **cache miss**
 - read block from memory (or lower level cache)
 - (evict an existing cache line to make room)
 - place new block in cache
 - return requested data
 - and stall the pipeline while all of this happens

Definitions and Terms

- Block (or line)
 - Minimum unit of information that is present/or not in the cache
- Cache hit, miss
- Hit rate
 - The fraction of memory accesses found in a level of the memory hierarchy
- Miss rate
 - The converse

Cache Questions

- What structure to use?
 - Where to place a block (book)?
 - How to find a block (book)?
- When miss, which block to replace?
- What happens on write?

Three common designs

A given data block can be placed...

- ... in exactly one cache line → Direct Mapped
- ... in any cache line → Fully Associative
- ... in a small set of cache lines → Set Associative

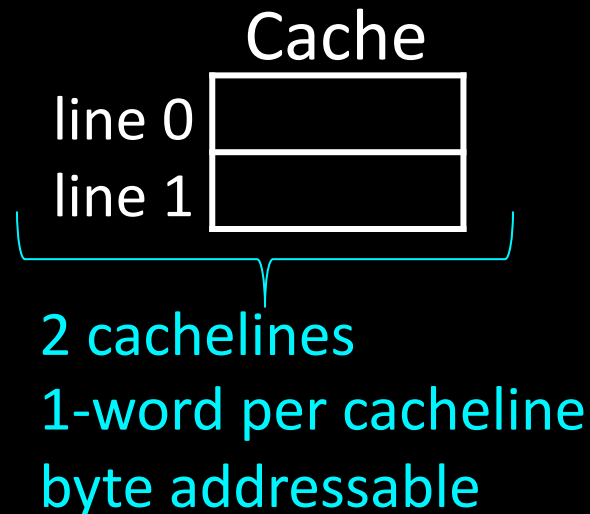
Direct Mapped Cache

- Each block number maps to a single cache line index
- Simplest hardware
- Questions
 - How to index into cache
 - How to find correct word/byte
 - How to match it

	Memory
0x000000	
0x000004	
0x000008	
0x00000c	
0x000010	
0x000014	
0x000018	
0x00001c	
0x000020	
0x000024	
0x000028	
0x00002c	
0x000030	
0x000034	
0x000038	
0x00003c	
0x000040	

Direct Mapped Cache

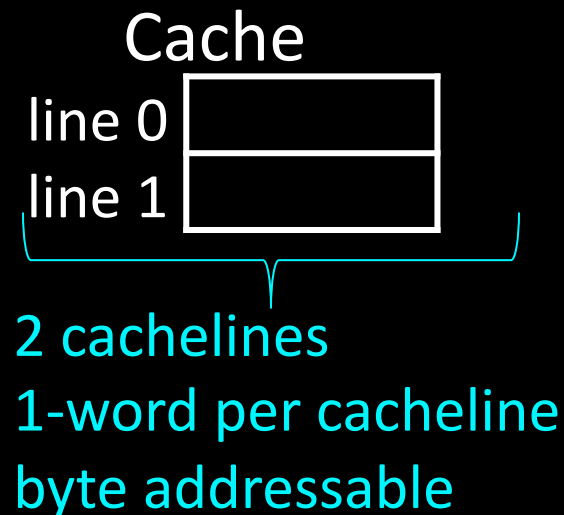
- Each block number maps to a single cache line index
- Simplest hardware
- Questions
 - How to index into cache
 - How to find correct word/byte
 - How to match it



Memory	
0x000000	
0x000004	
0x000008	
0x00000c	
0x000010	
0x000014	
0x000018	
0x00001c	
0x000020	
0x000024	
0x000028	
0x00002c	
0x000030	
0x000034	
0x000038	
0x00003c	
0x000040	

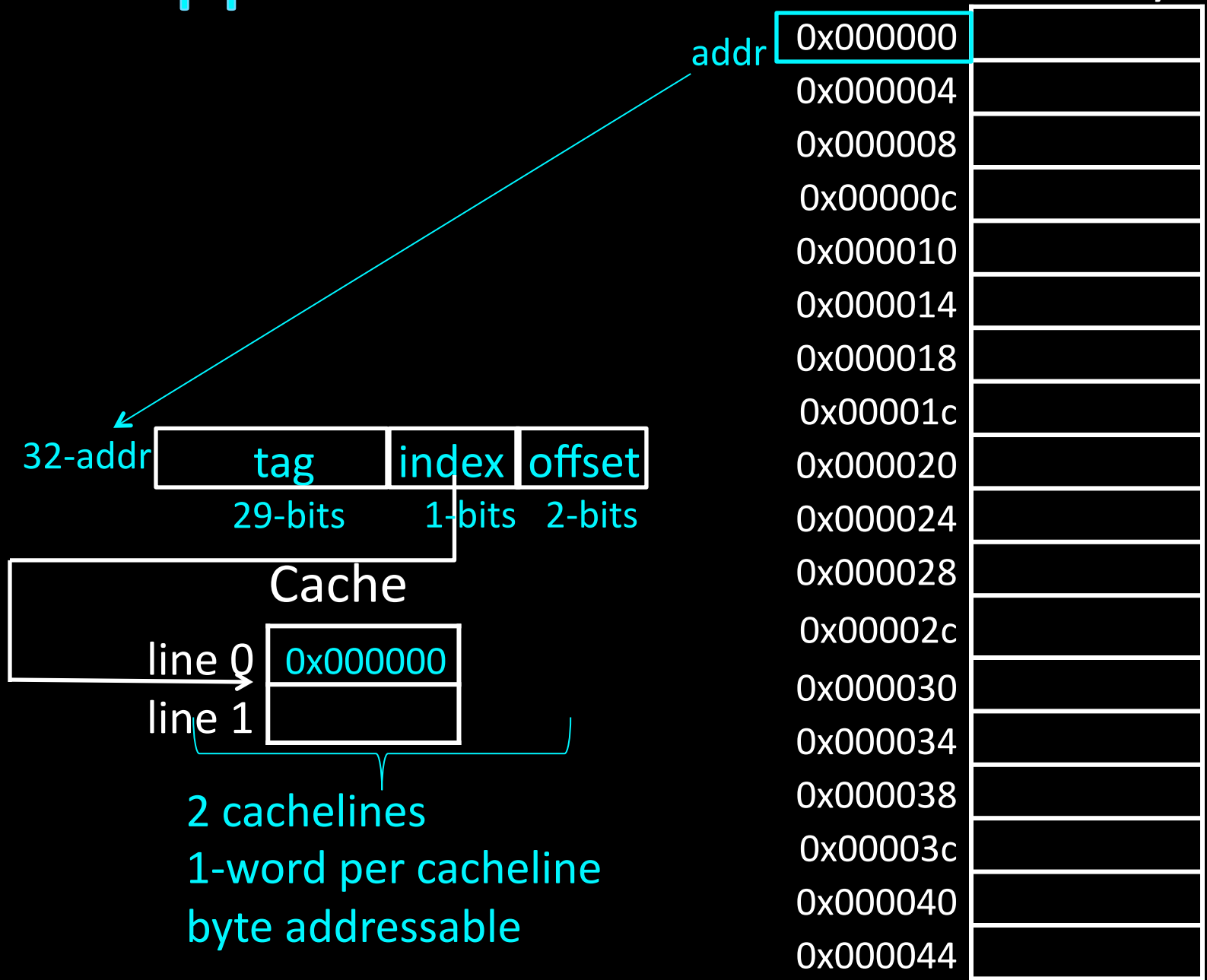
Direct Mapped Cache

- Questions
 - How to index into cache
 - How to find correct word/byte
 - How to match it



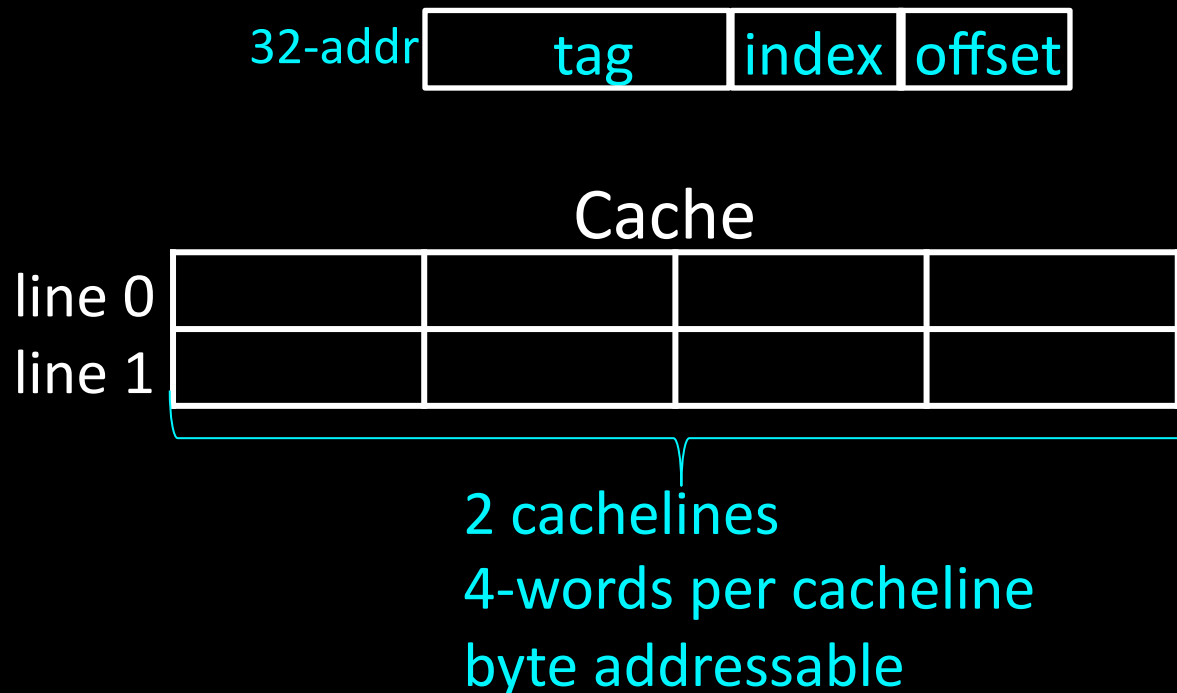
Memory	
0x000000	
0x000004	
0x000008	
0x00000c	
0x000010	
0x000014	
0x000018	
0x00001c	
0x000020	
0x000024	
0x000028	
0x00002c	
0x000030	
0x000034	
0x000038	
0x00003c	
0x000040	

Direct Mapped Cache



Direct Mapped Cache

- Each block number maps to a single cache line index
- Simplest hardware



Memory	
0x000000	
0x000004	
0x000008	
0x00000c	
0x000010	
0x000014	
0x000018	
0x00001c	
0x000020	
0x000024	
0x000028	
0x00002c	
0x000030	
0x000034	
0x000038	
0x00003c	
0x000040	

iClicker

Size of offset?

A) 1

B) 2

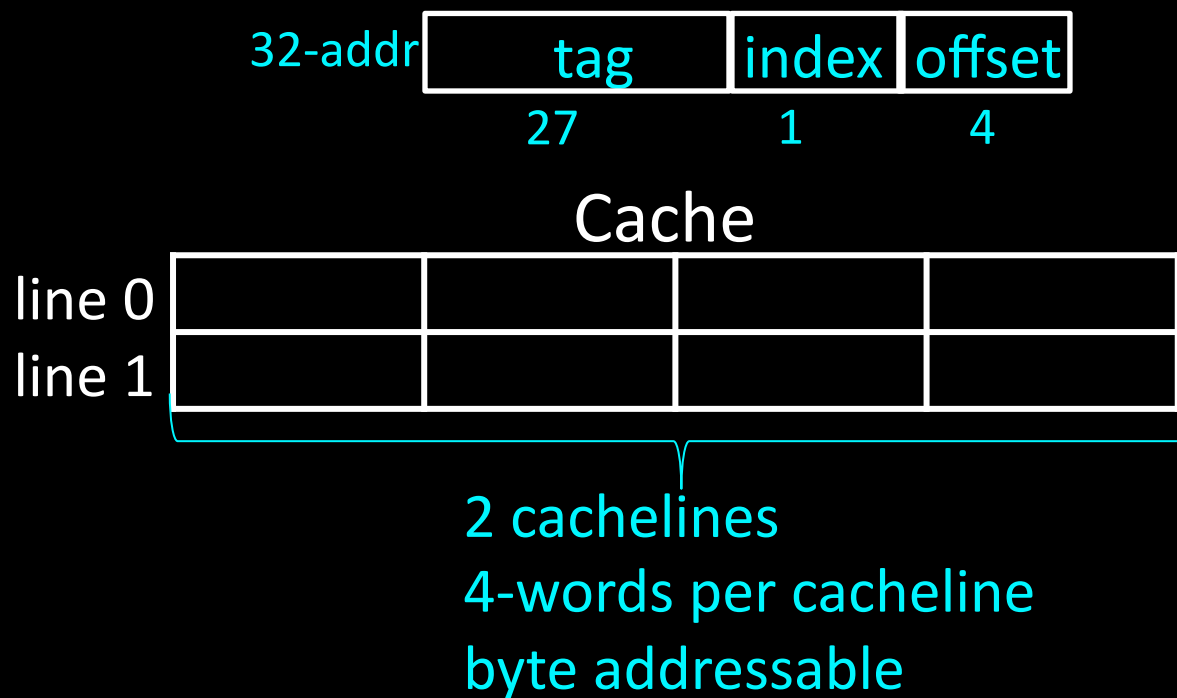
C) 3

D) 4

E) 5

Size of tag?

Direct Mapped Cache



Memory

0x000000

0x000004

0x000008

0x00000c

0x000010

0x000014

0x000018

0x00001c

0x000020

0x000024

0x000028

0x00002c

0x000030

0x000034

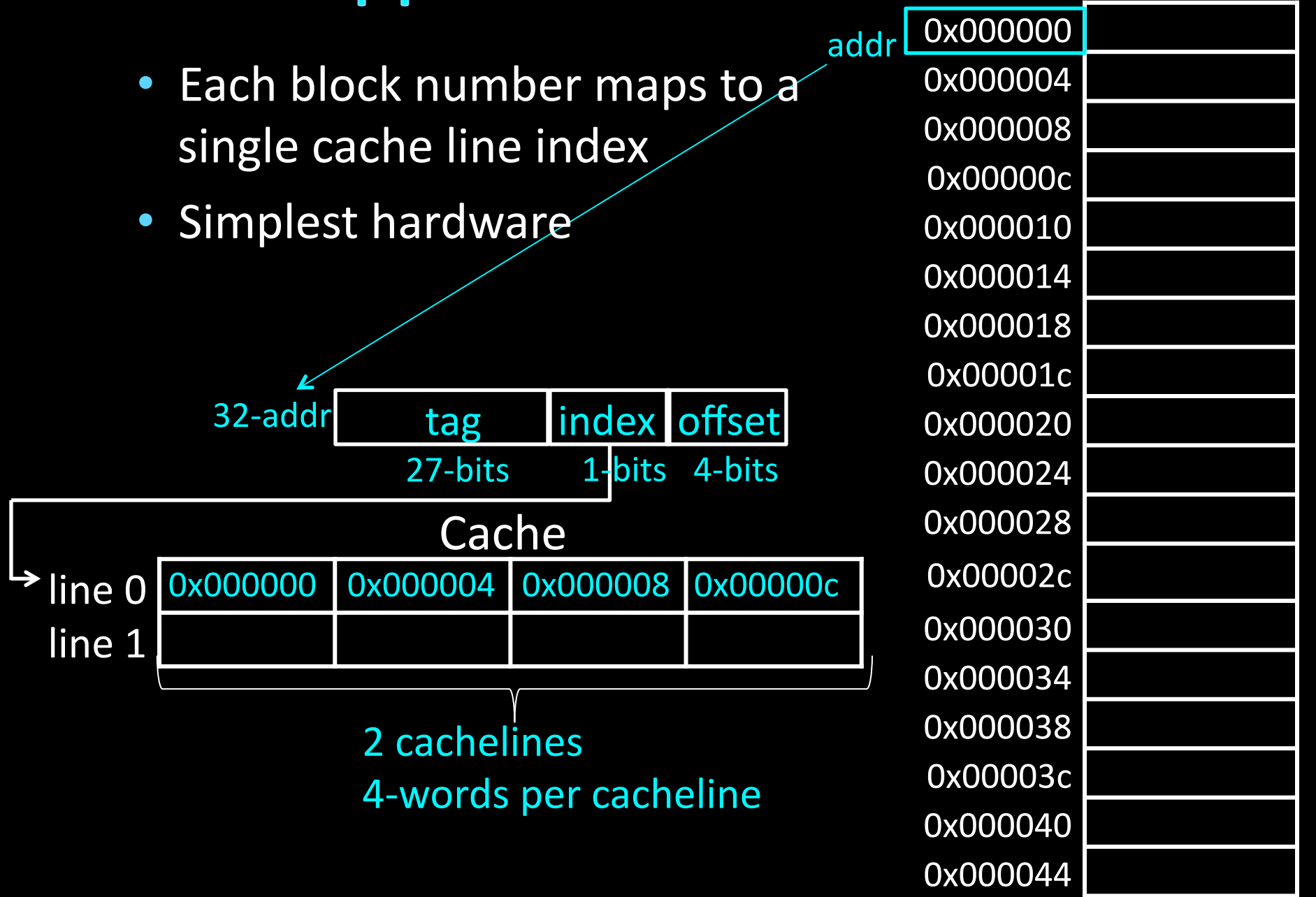
0x000038

0x00003c

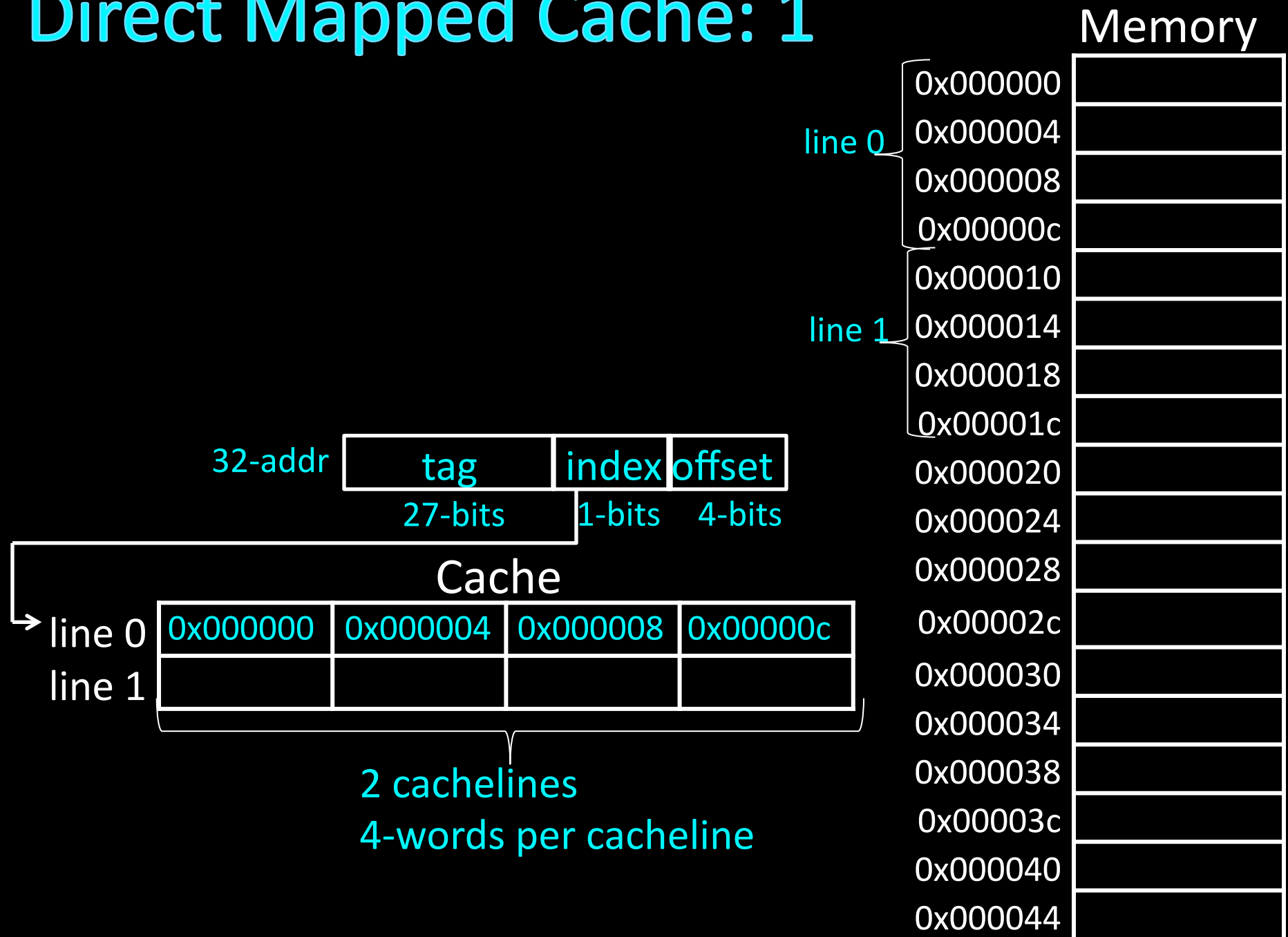
0x000040

Direct Mapped Cache

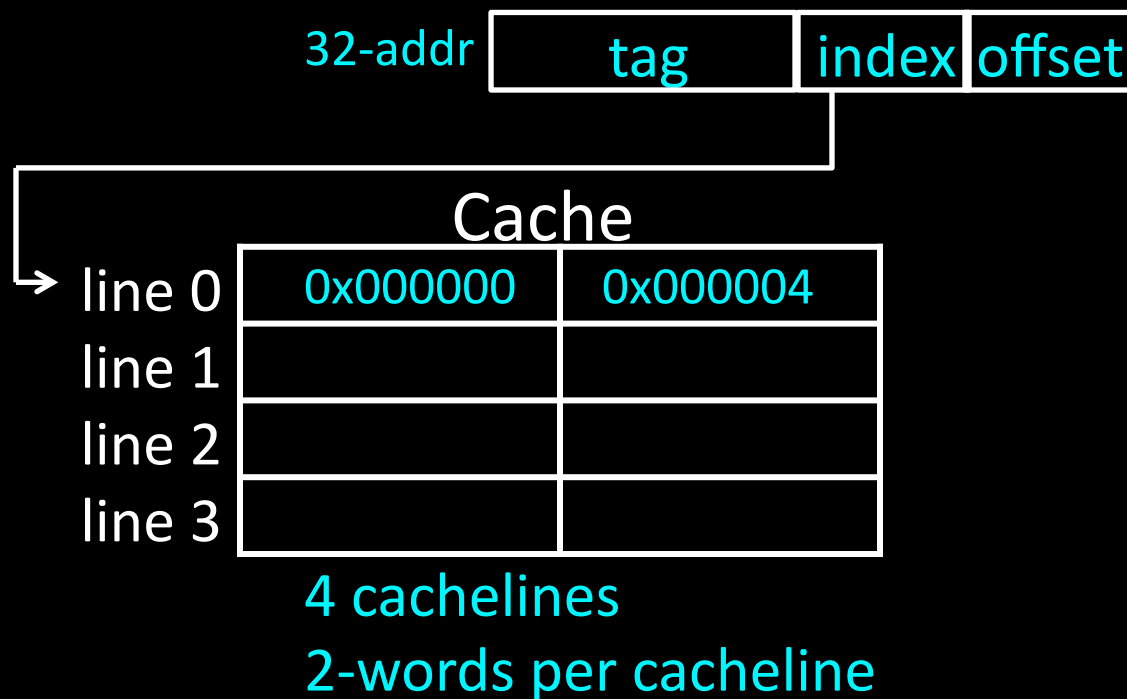
- Each block number maps to a single cache line index
- Simplest hardware



Direct Mapped Cache: 1

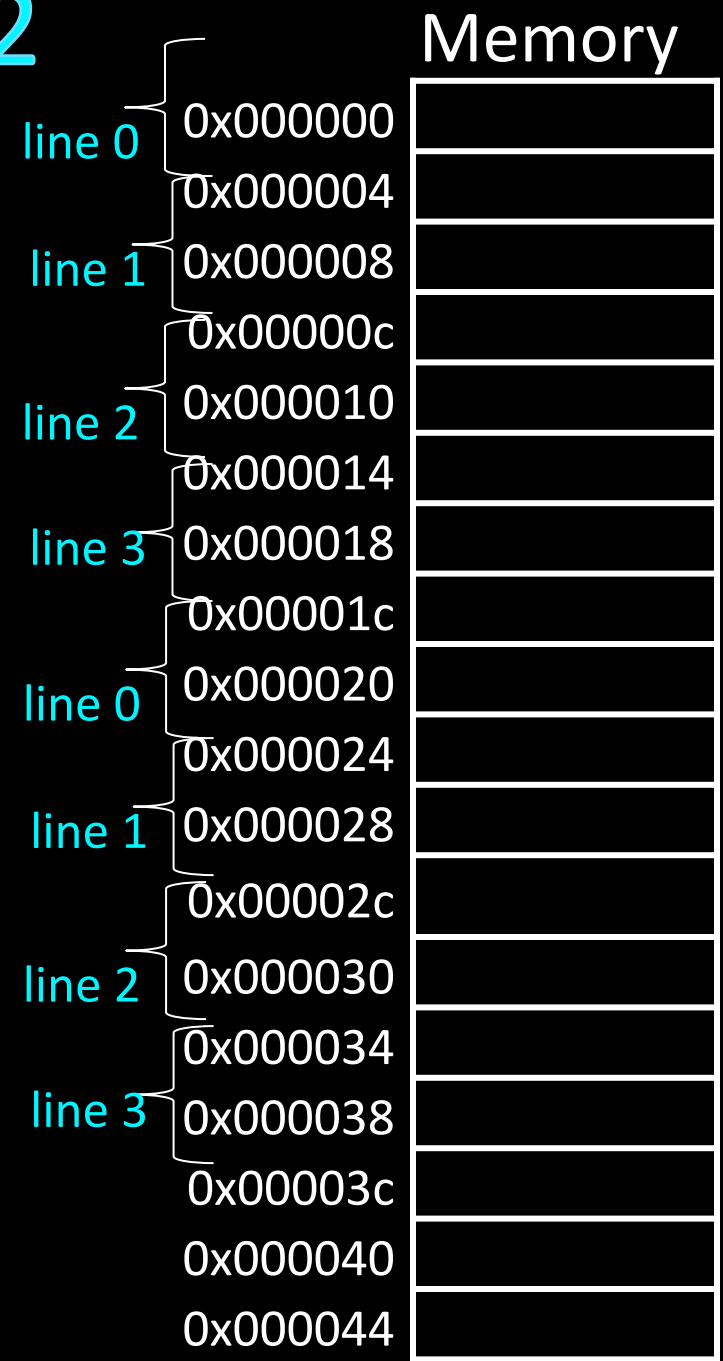
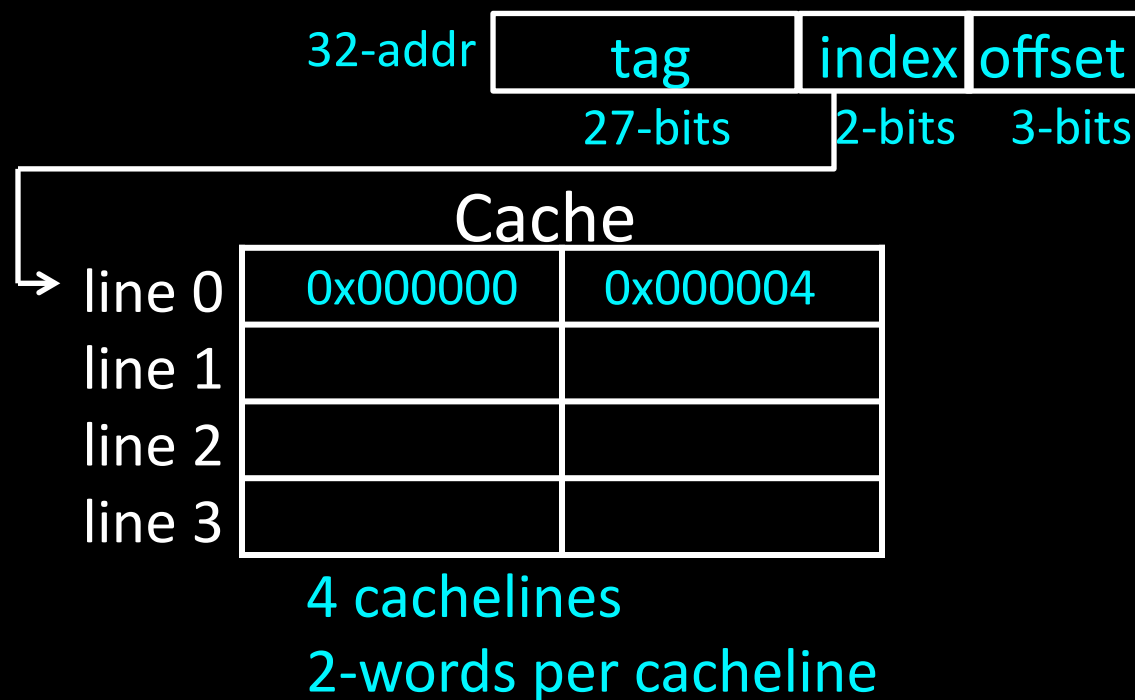


Direct Mapped Cache: 2



	Memory
0x000000	
0x000004	
0x000008	
0x00000c	
0x000010	
0x000014	
0x000018	
0x00001c	
0x000020	
0x000024	
0x000028	
0x00002c	
0x000030	
0x000034	
0x000038	
0x00003c	
0x000040	
0x000044	

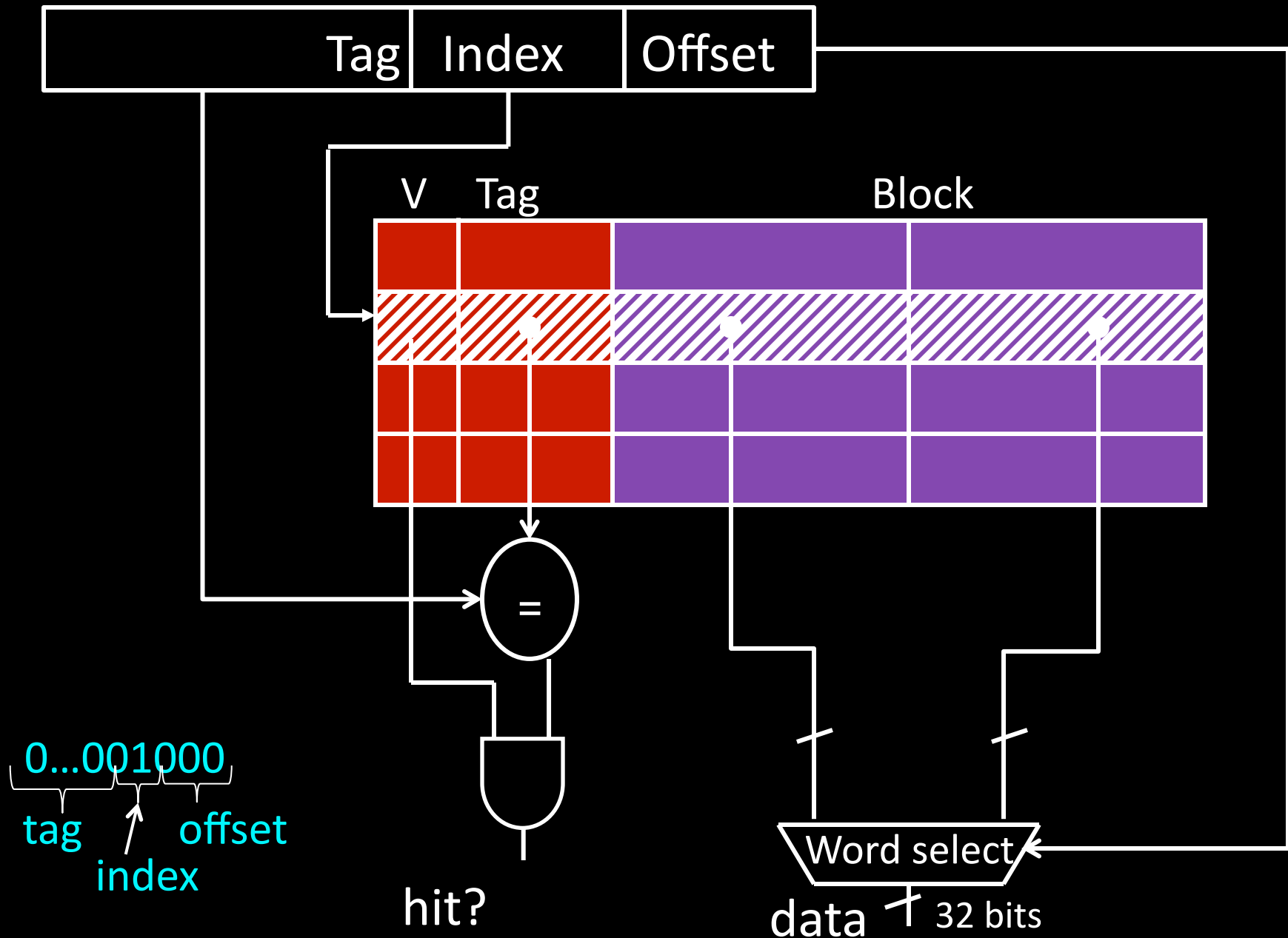
Direct Mapped Cache: 2



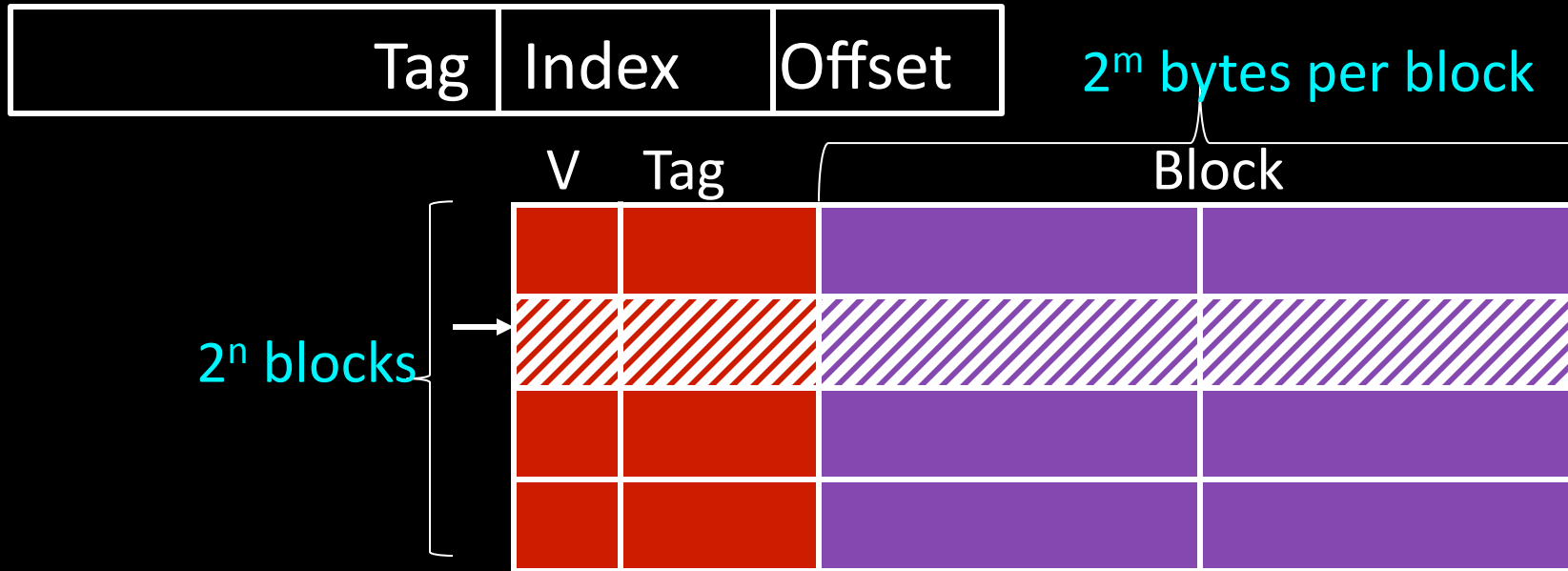
Direct Mapped Cache

Pros: Very simple hardware

Direct Mapped Cache (Hardware)



Direct Mapped Cache Size



n bit index, m bit offset

Q: How big is cache (**data only**)?

Cache of size 2^n blocks

Block size of 2^m bytes

Cache Size: 2^n bytes per block x 2^n blocks = 2^{n+m} bytes

Direct Mapped Cache Size



n bit index, m bit offset

Q: How much SRAM is needed (**data + overhead**)?

Cache of size 2^n blocks

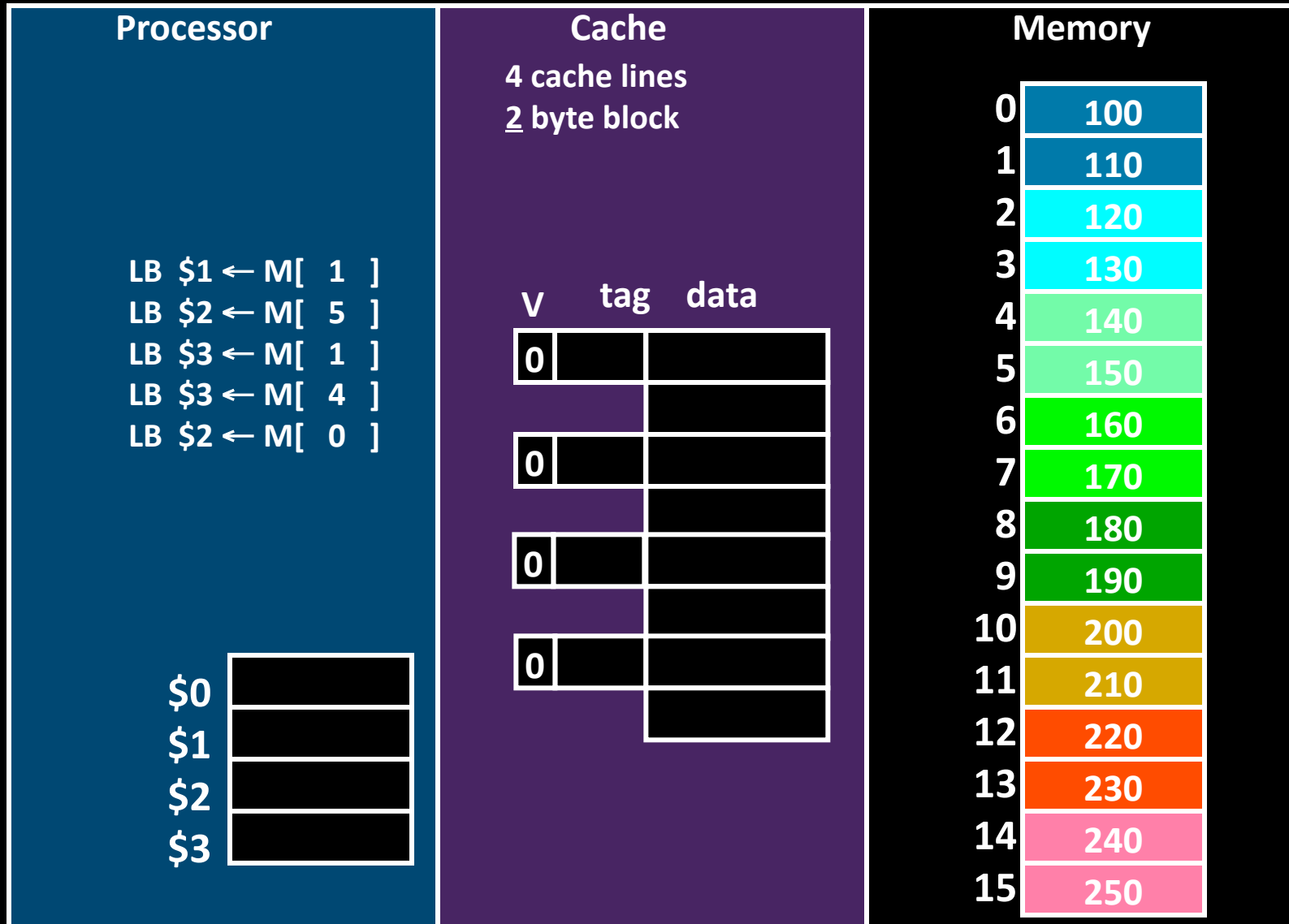
Block size of 2^m bytes

Tag field: $32 - (n + m)$, Valid bit: 1

SRAM Size: $2^n \times (\text{block size} + \text{tag size} + \text{valid bit size})$
 $= 2^n \times (2^m \text{ bytes} \times 8 \text{ bits-per-byte} + (32 - n - m) + 1) \text{ bits}$

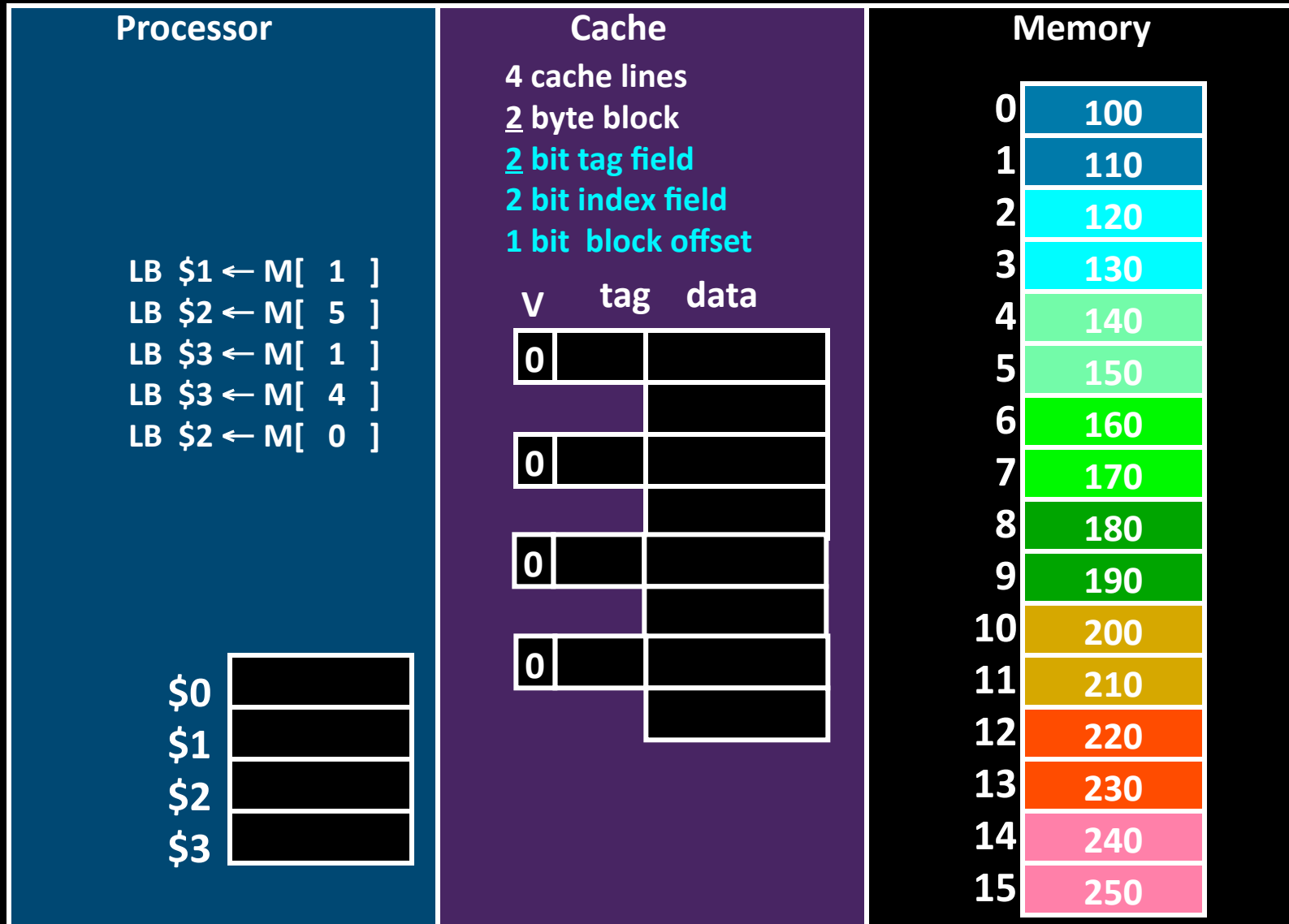
Example: Direct Mapped Cache

Using **byte addresses** in this example. Addr Bus = 5 bits

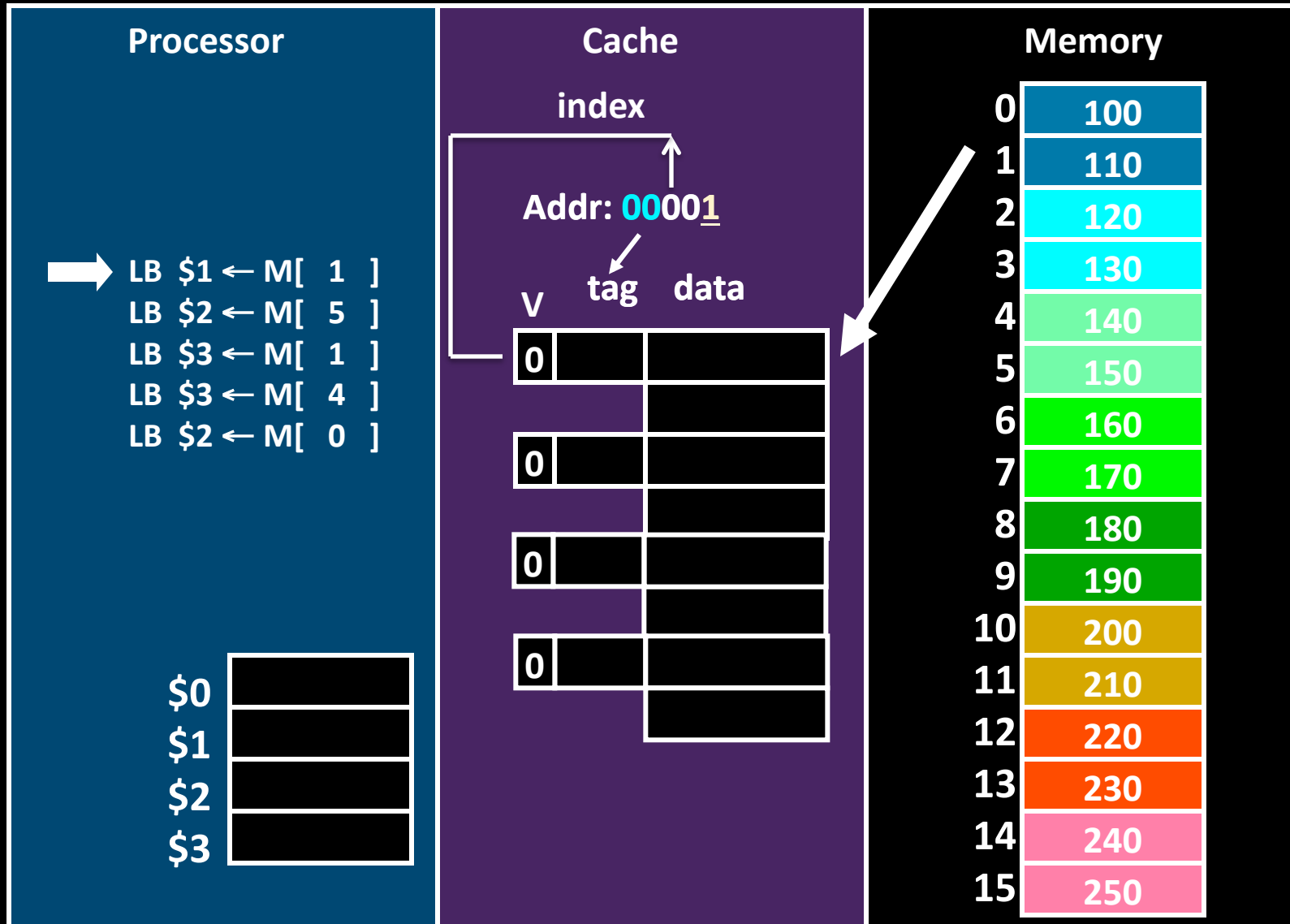


Example: Direct Mapped Cache

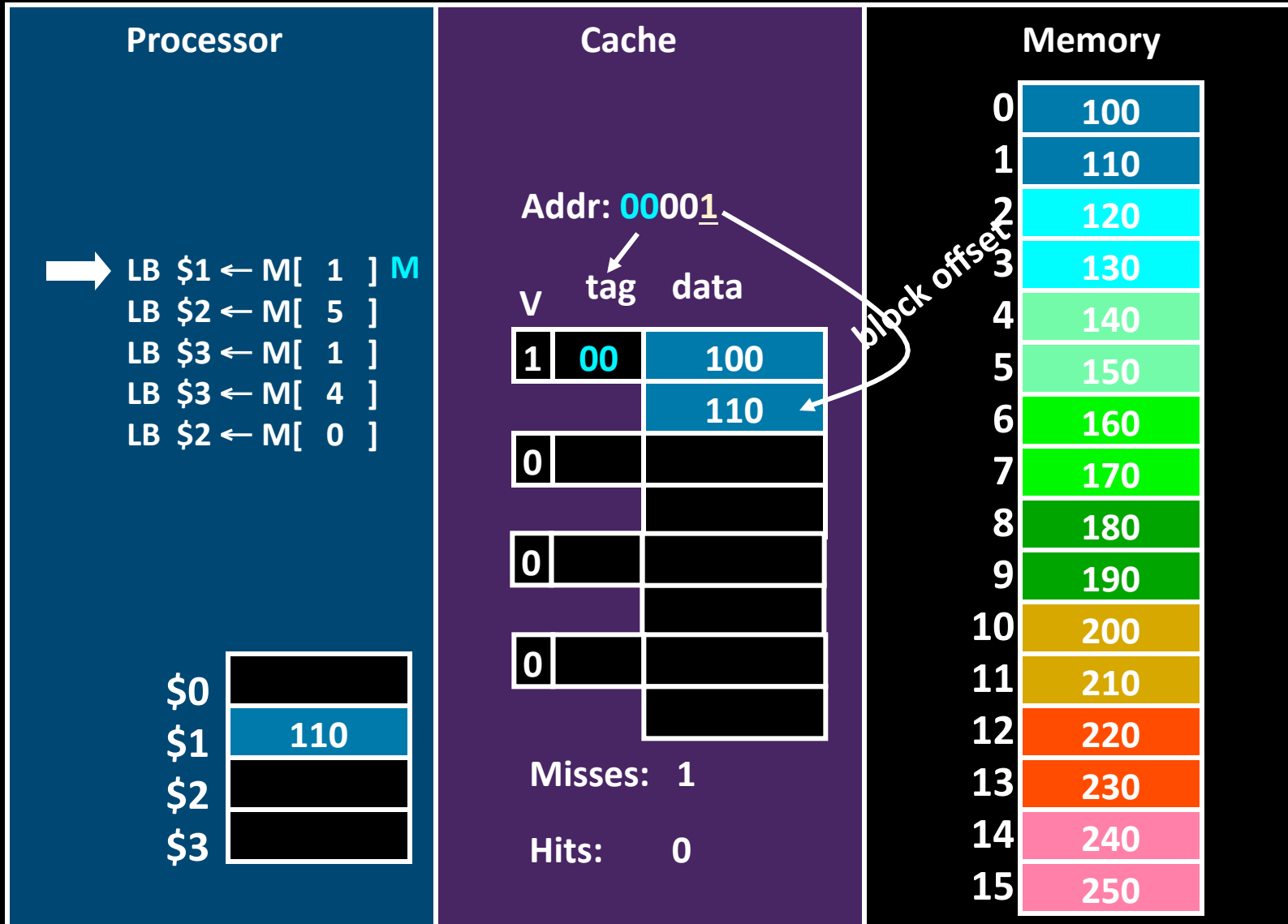
Using **byte addresses** in this example. Addr Bus = 5 bits



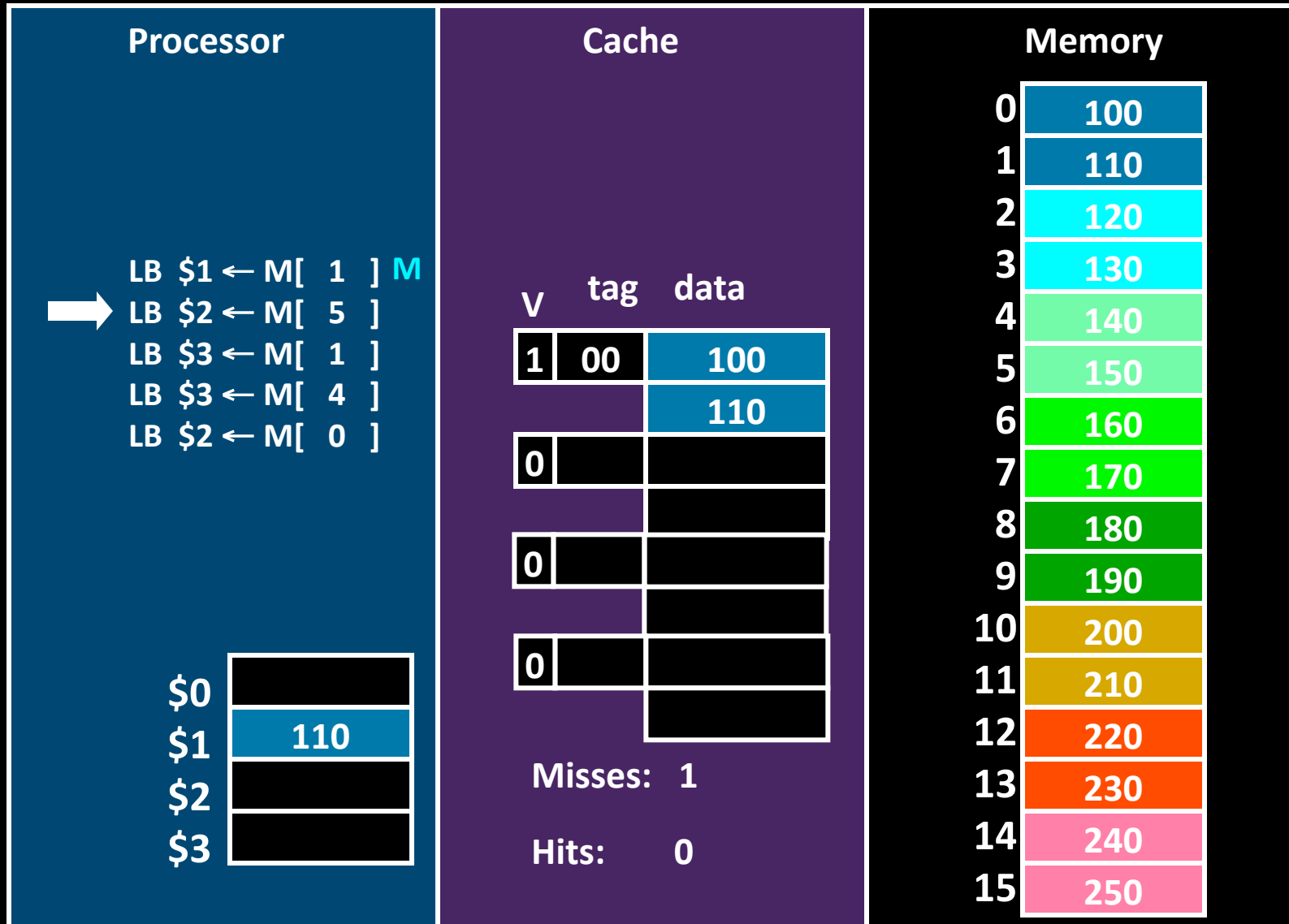
1st Access



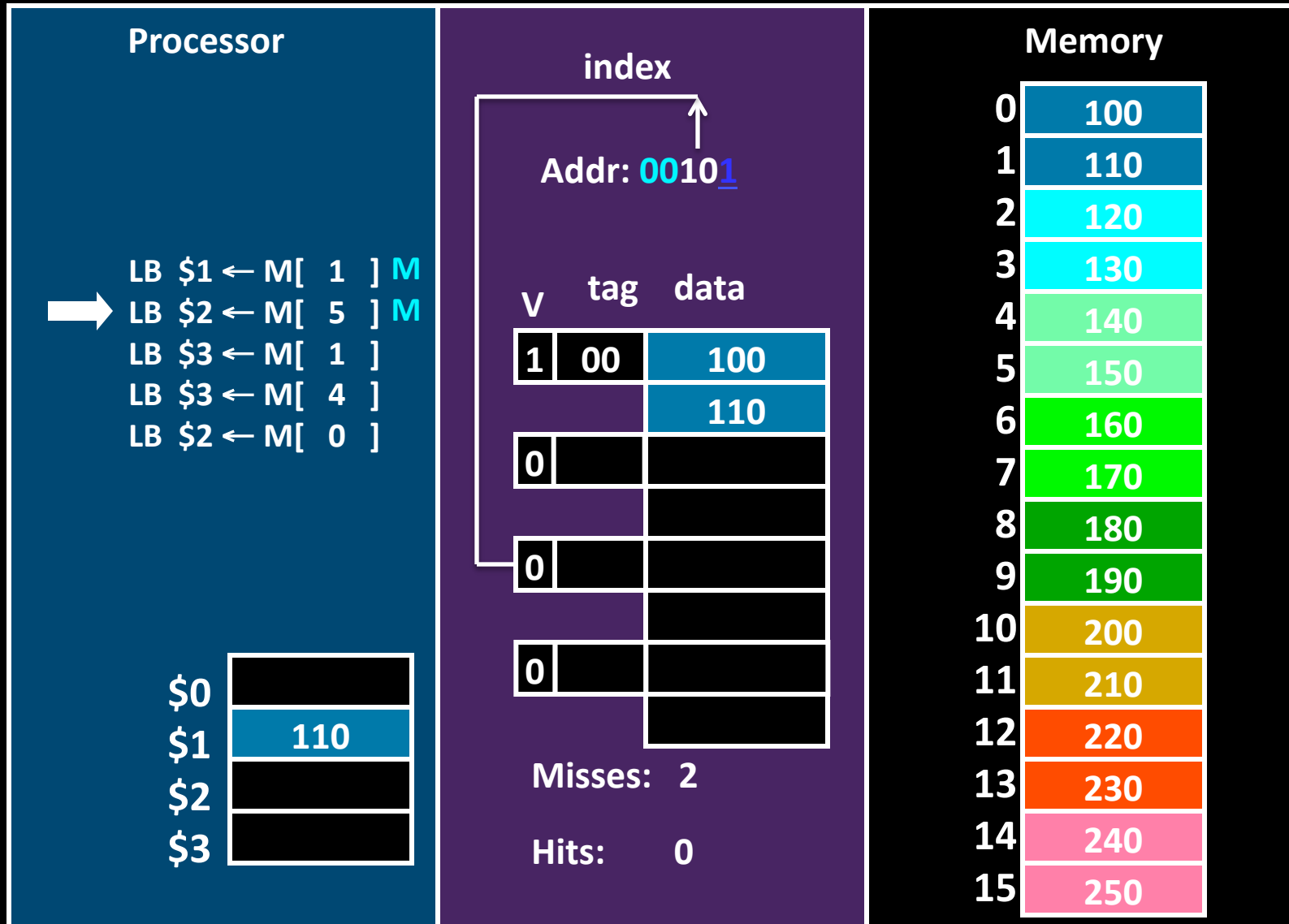
1st Access



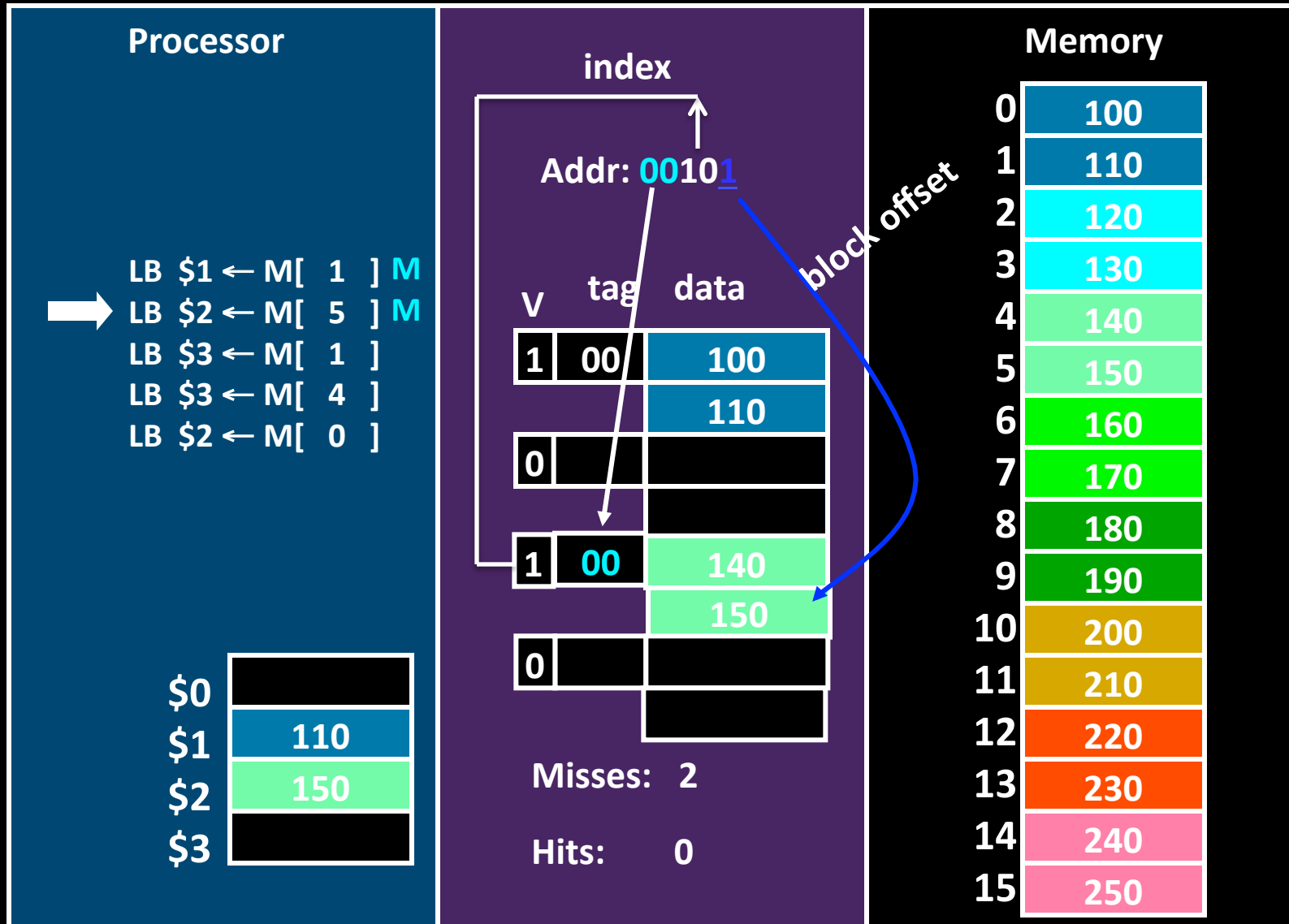
2nd Access



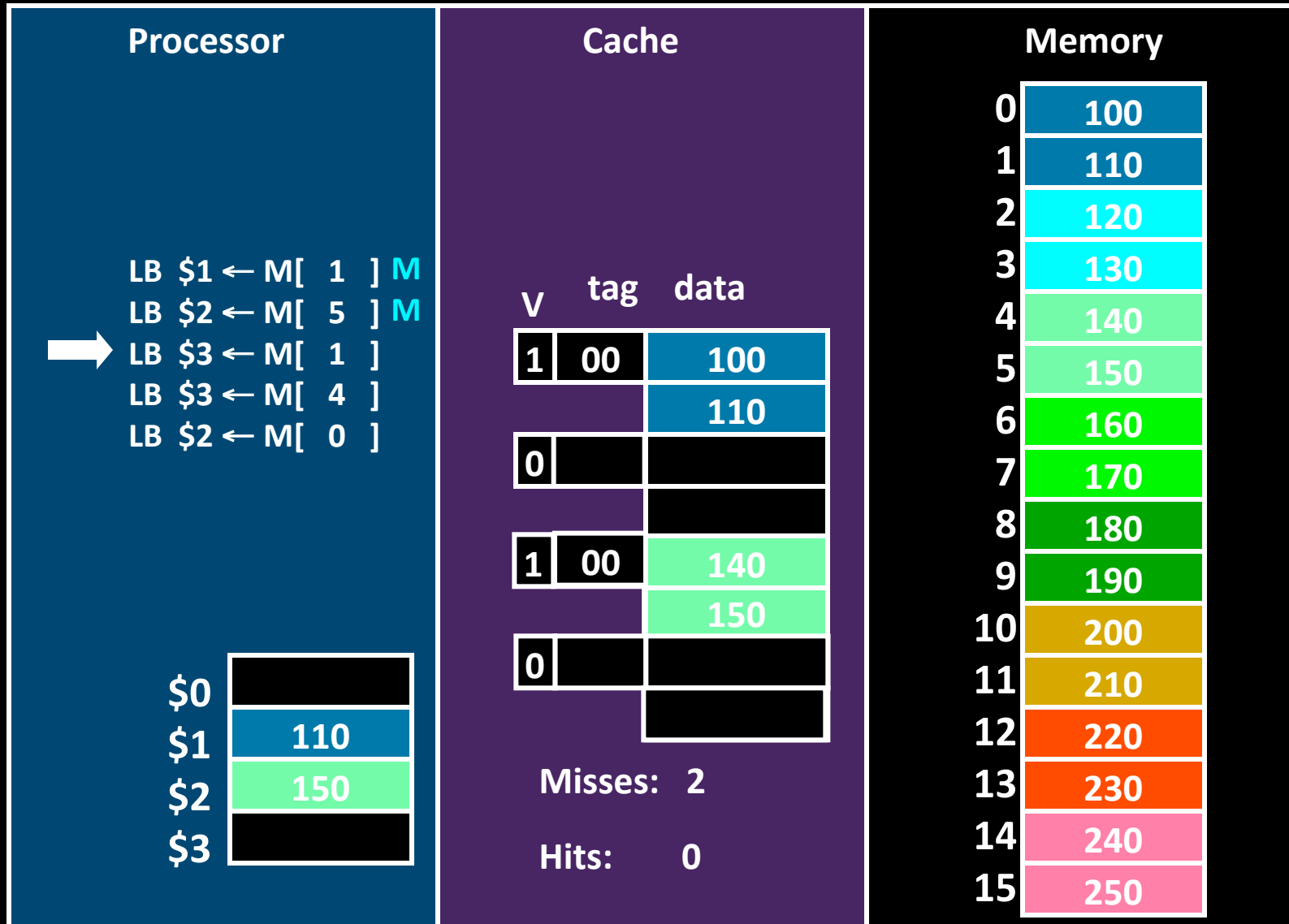
2nd Access



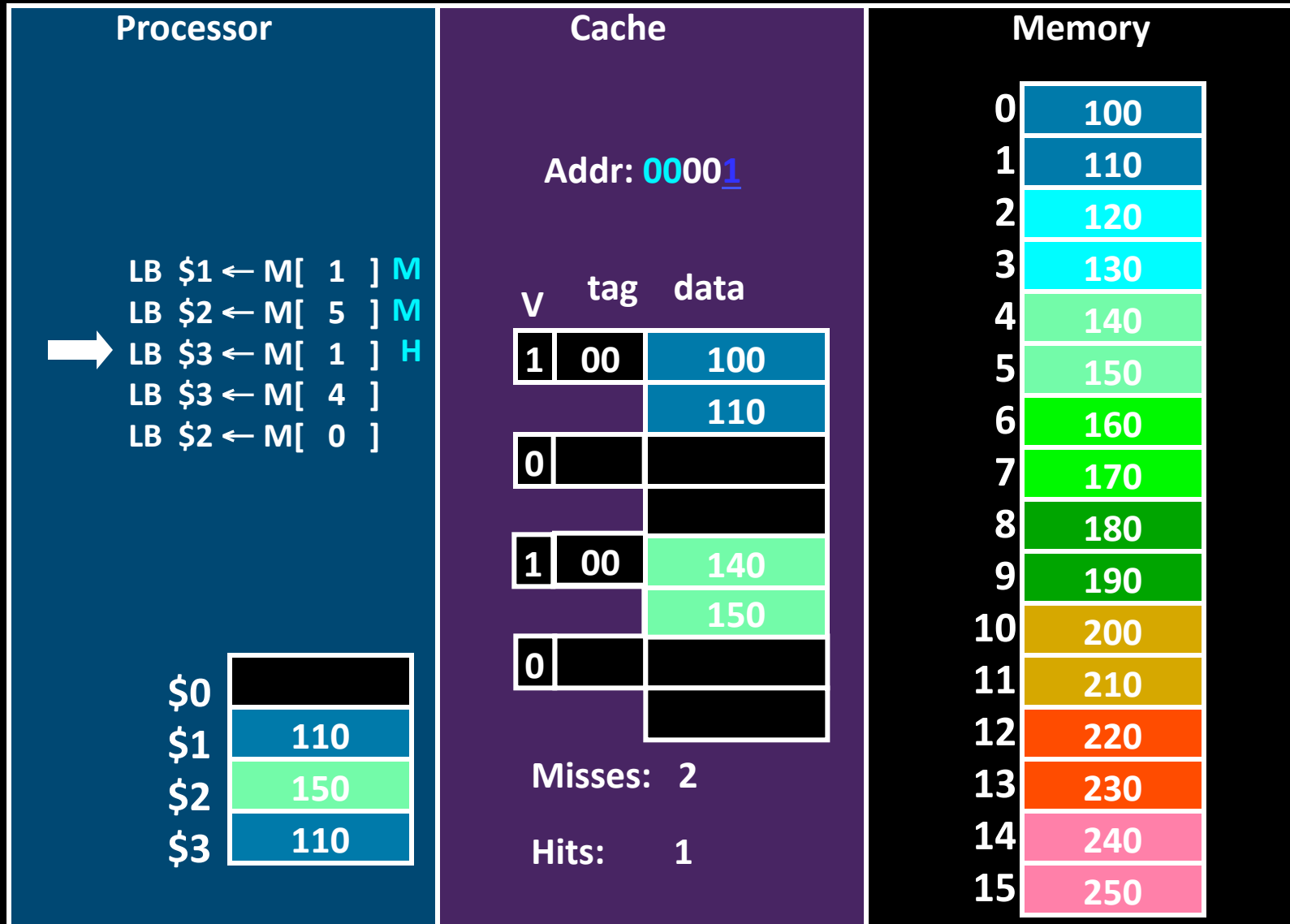
2nd Access



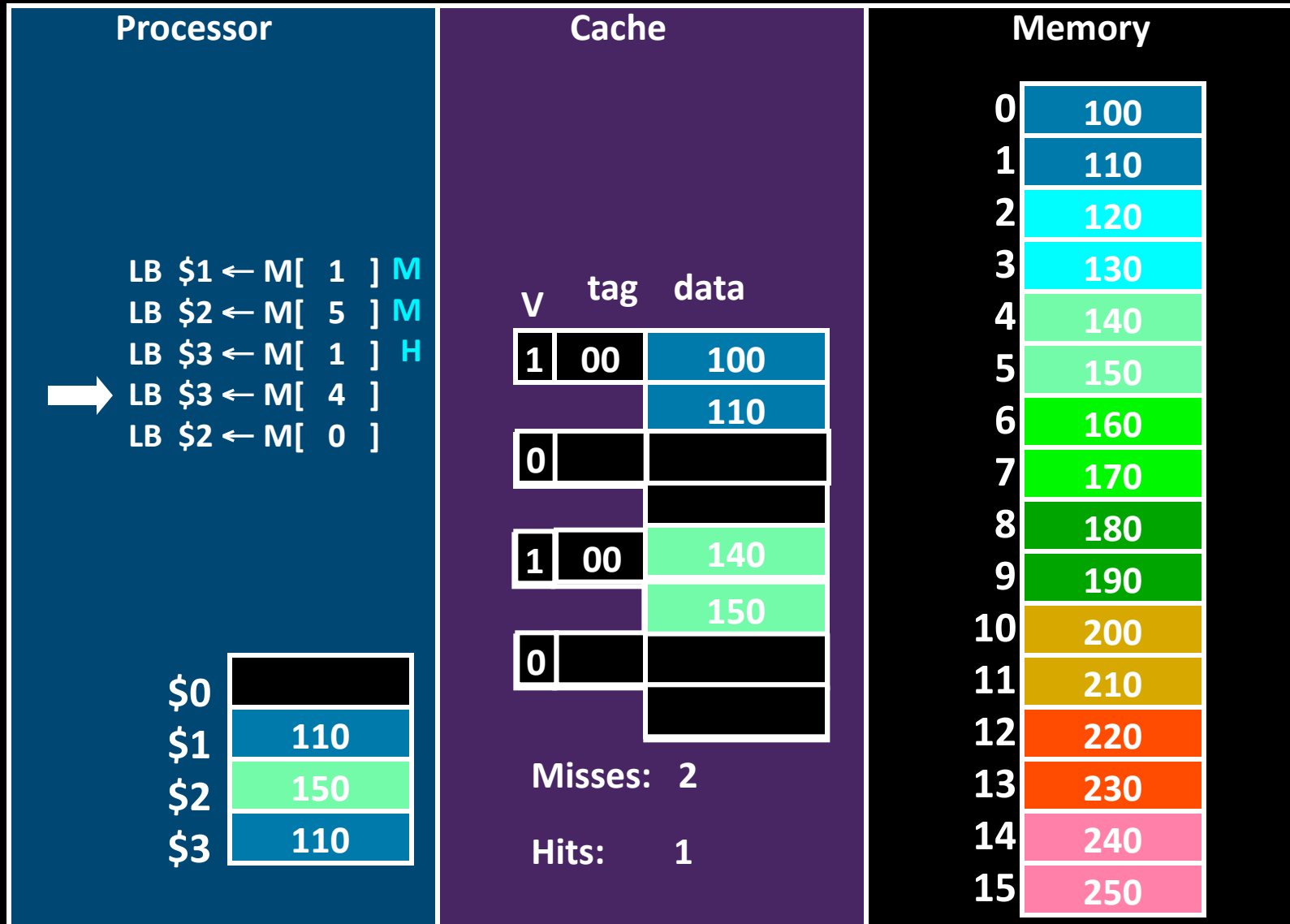
3rd Access



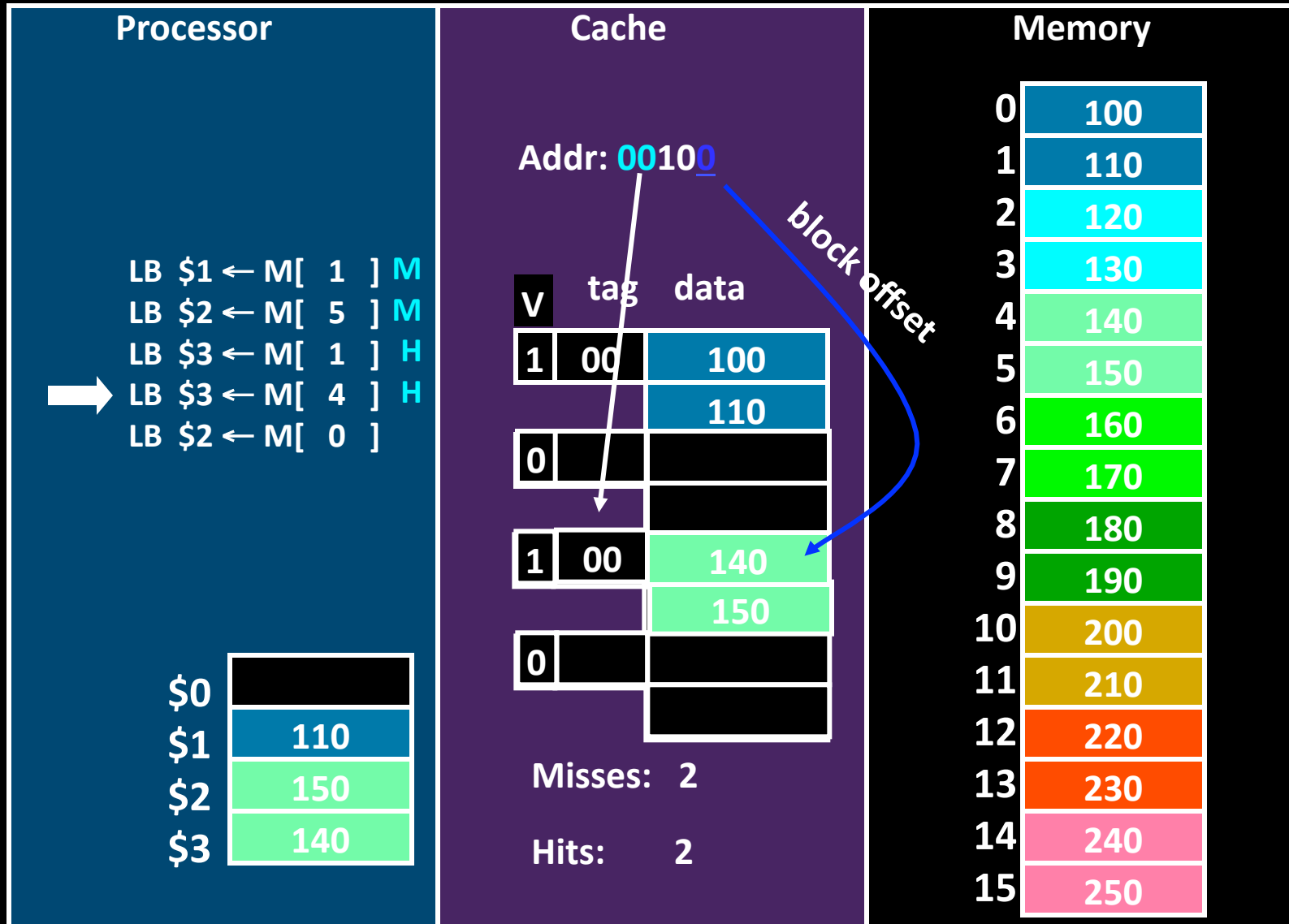
3rd Access



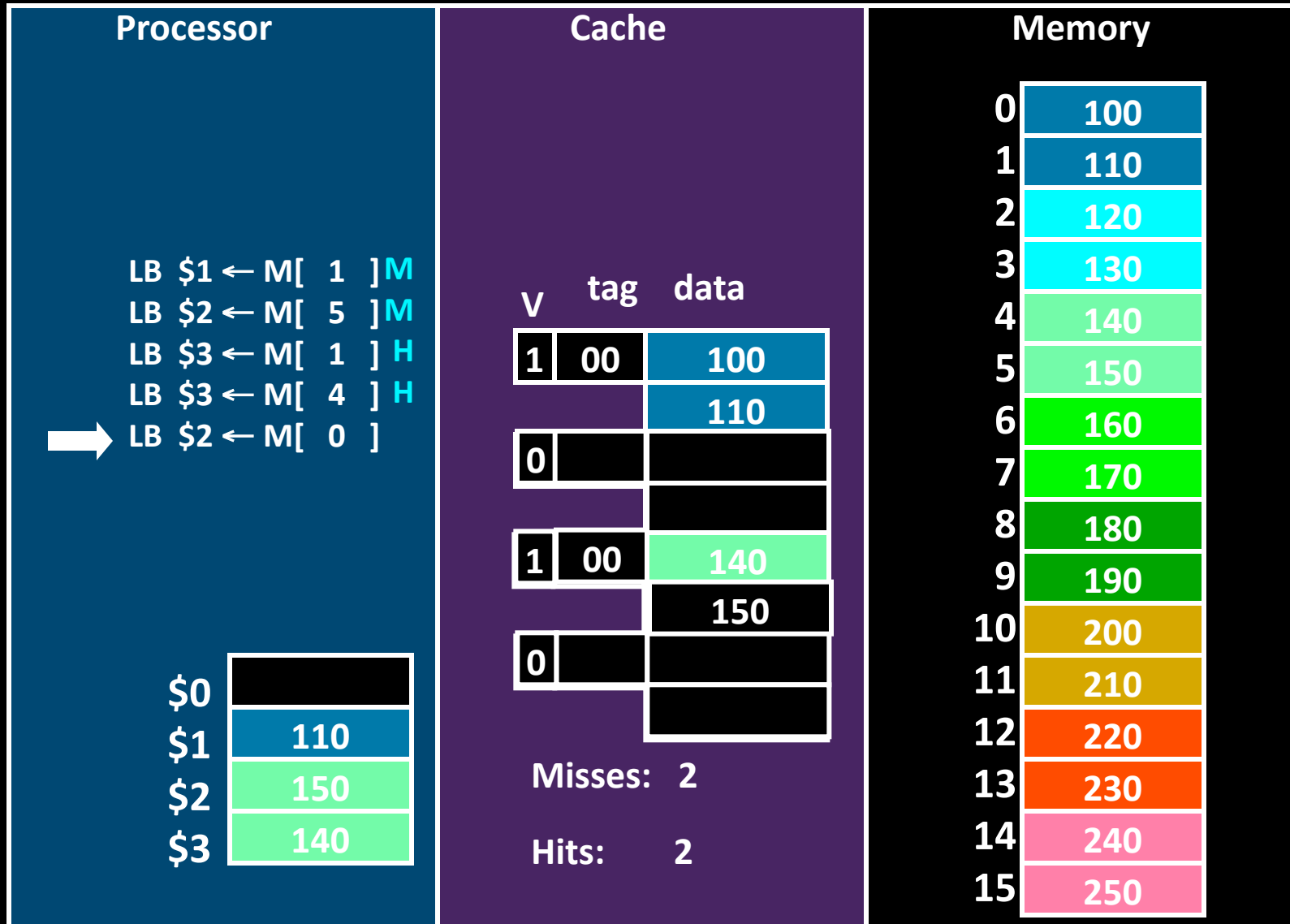
4th Access



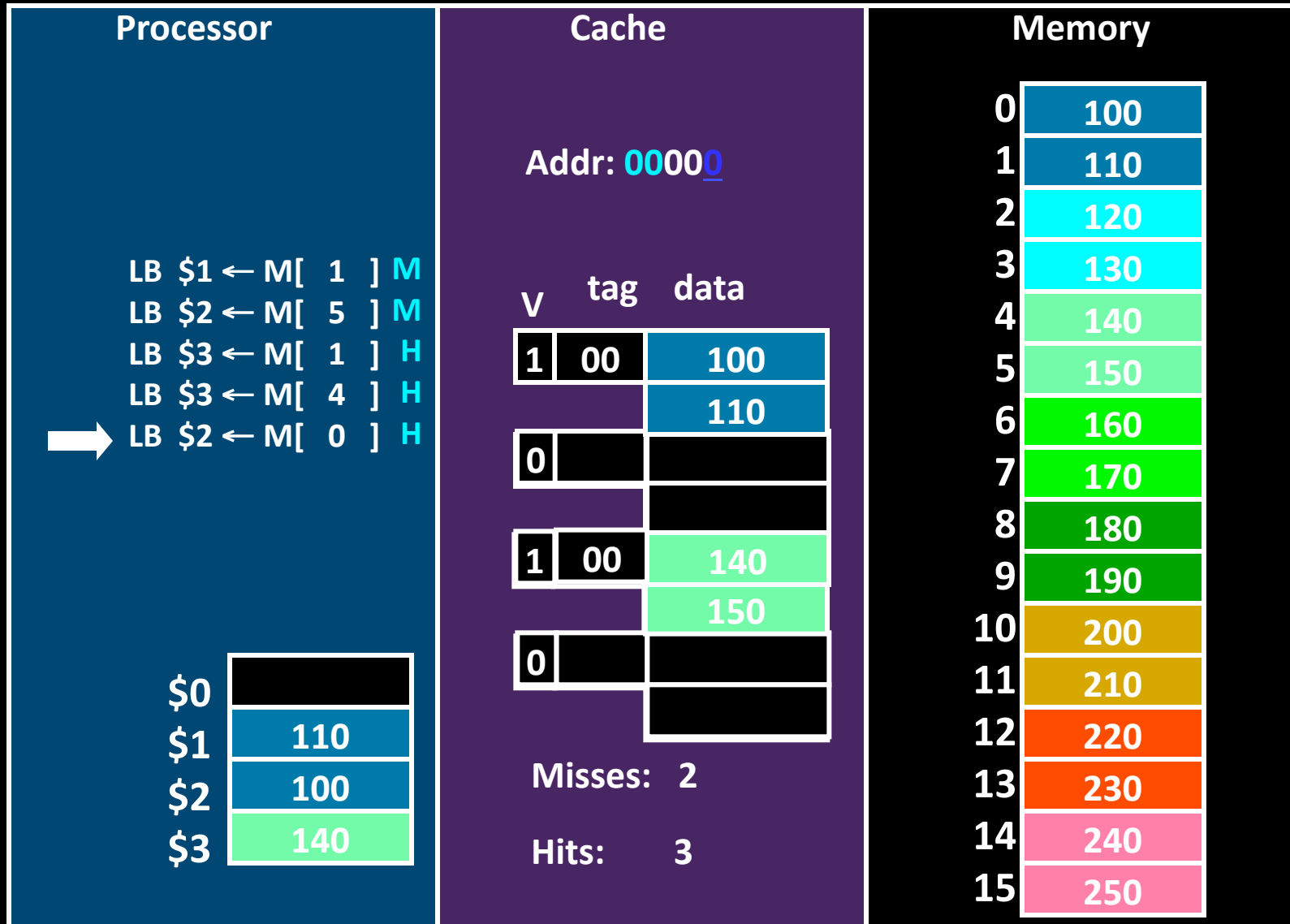
4th Access



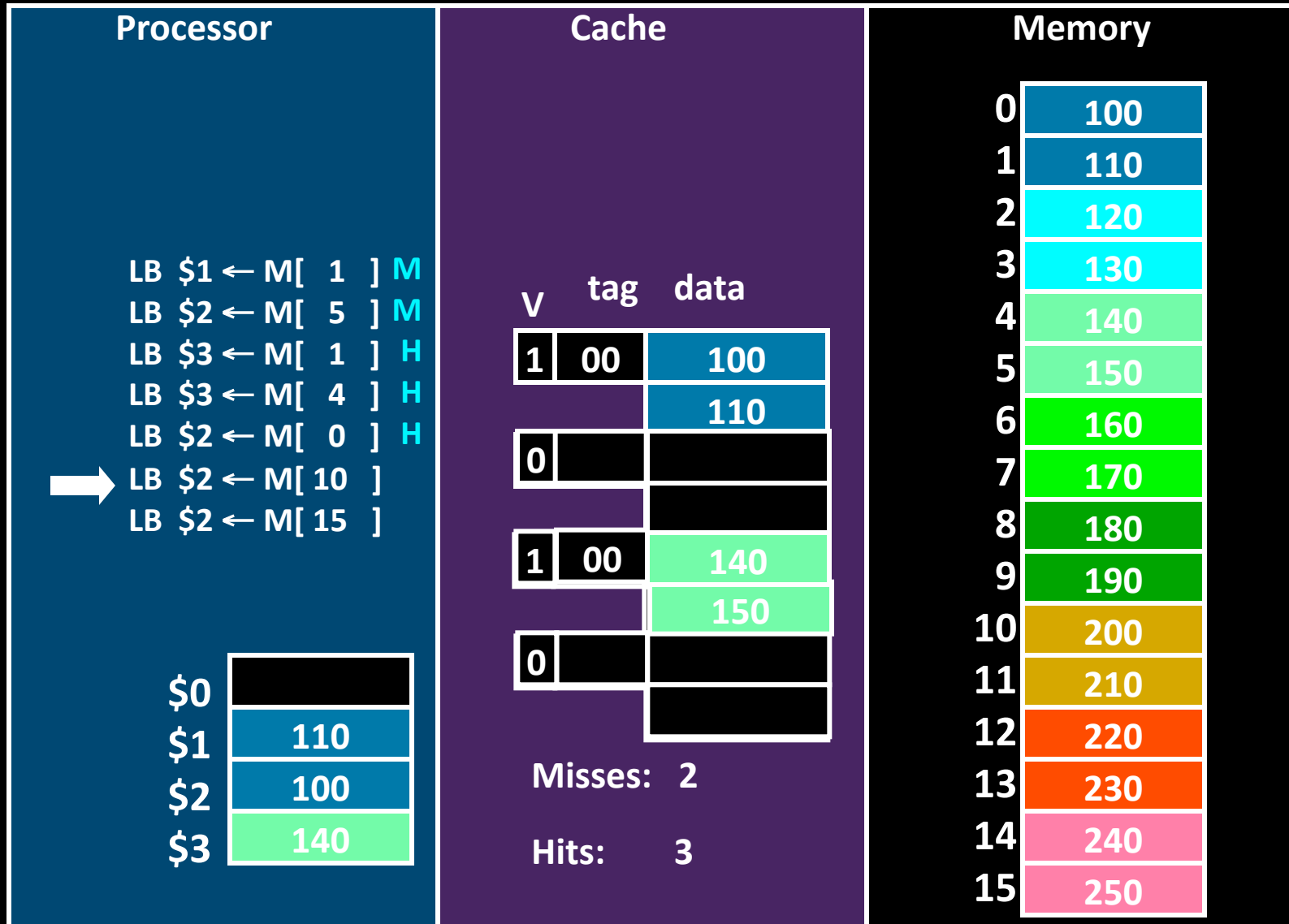
5th Access



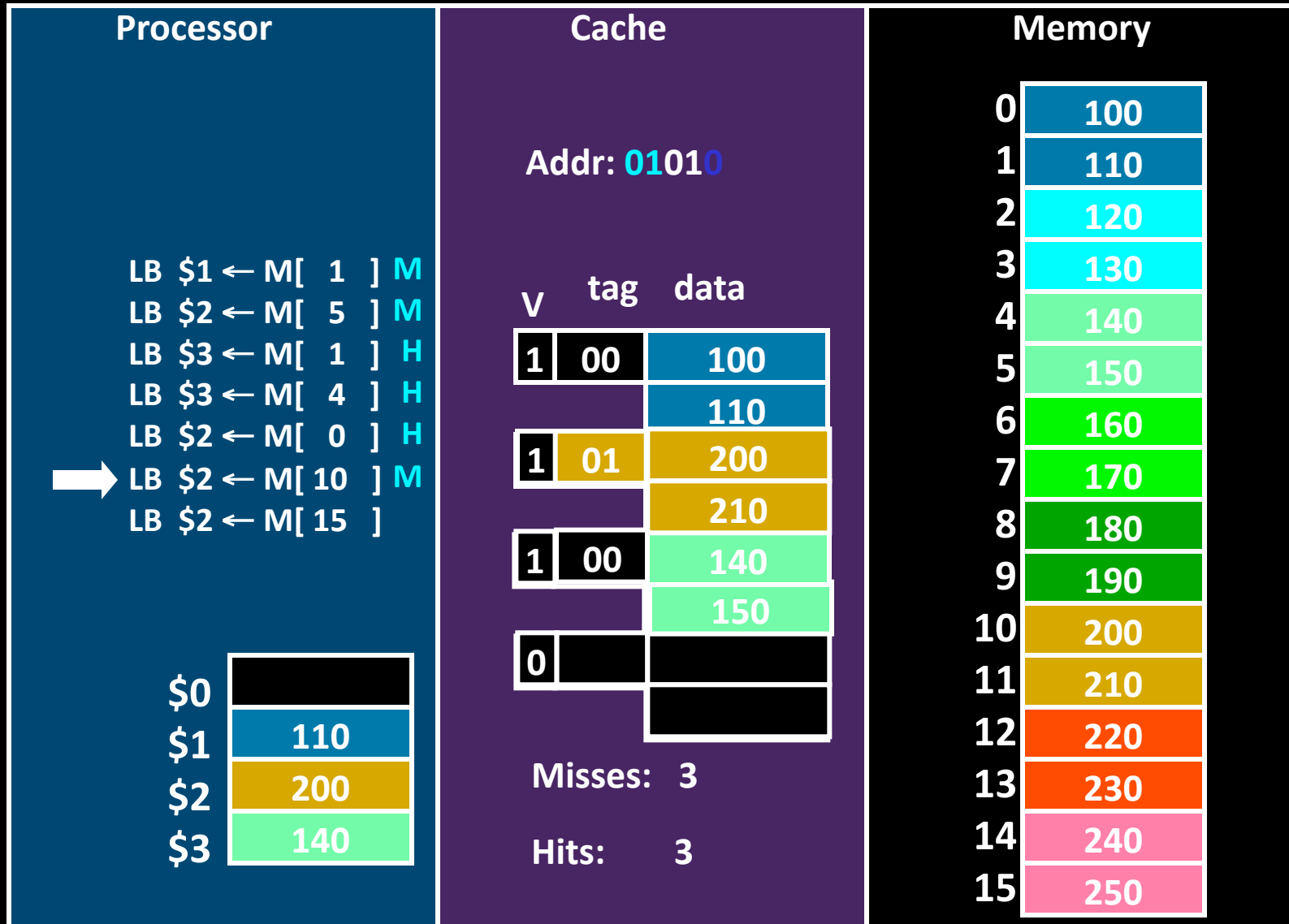
5th Access



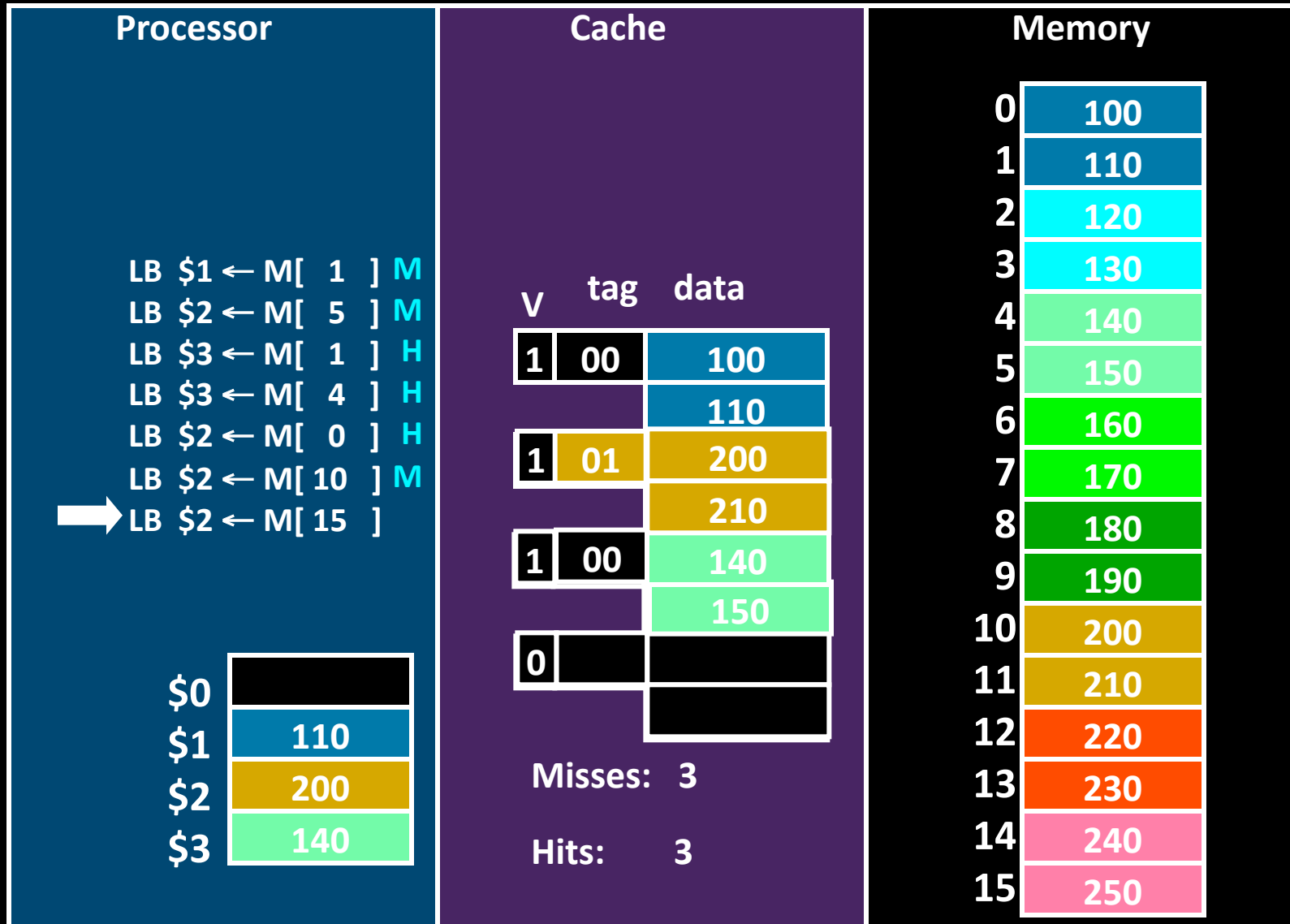
6th Access



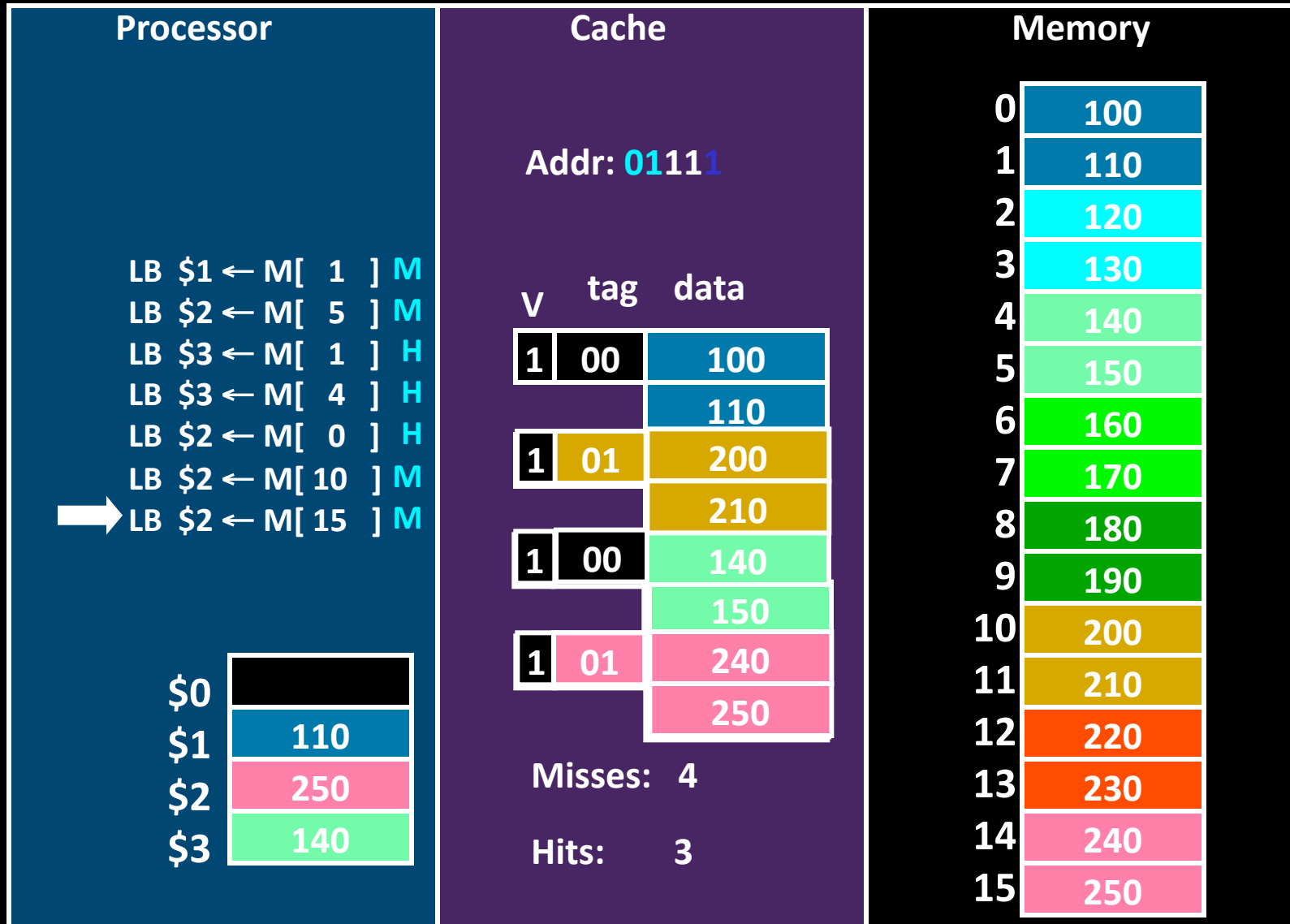
6th Access



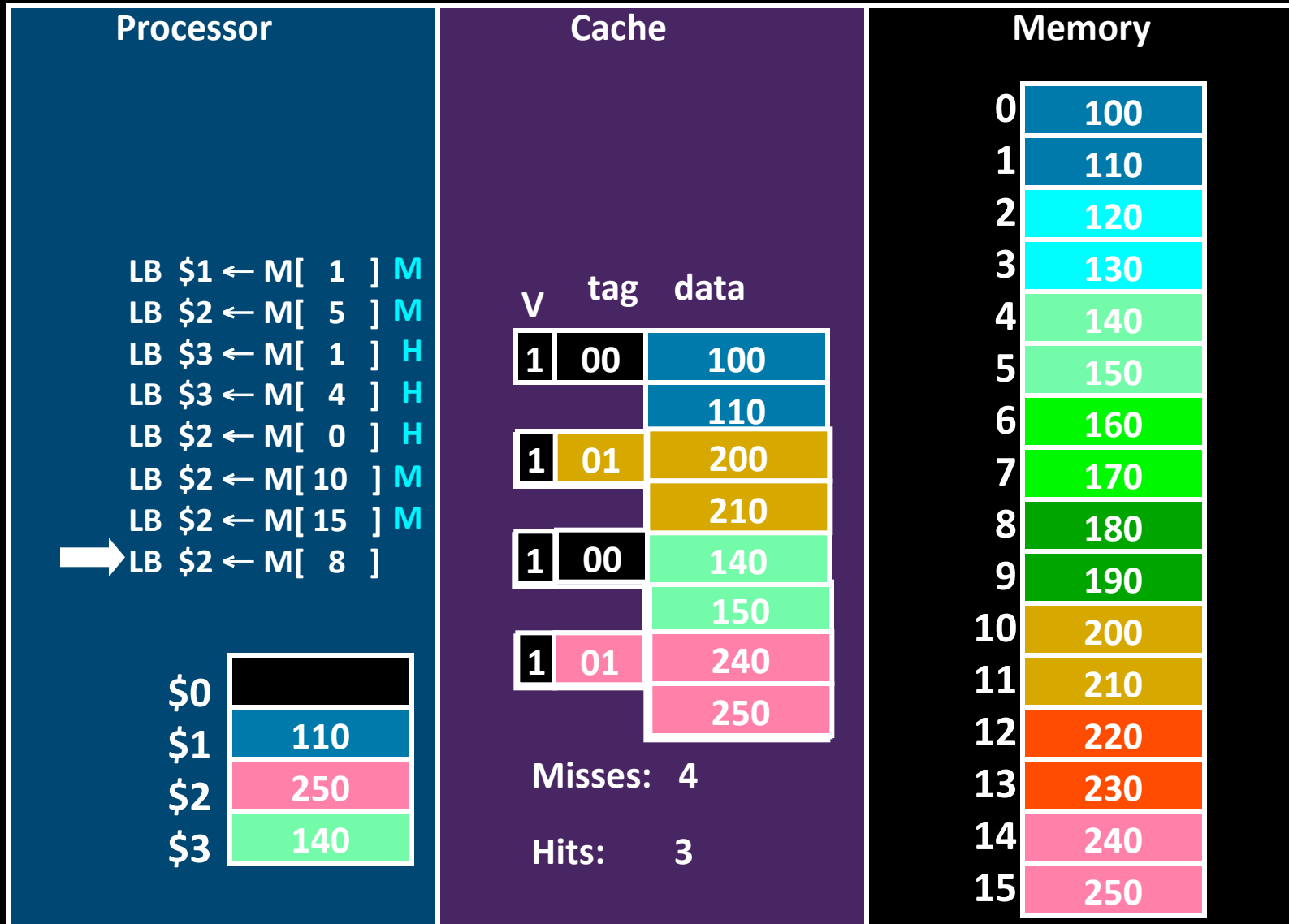
7th Access



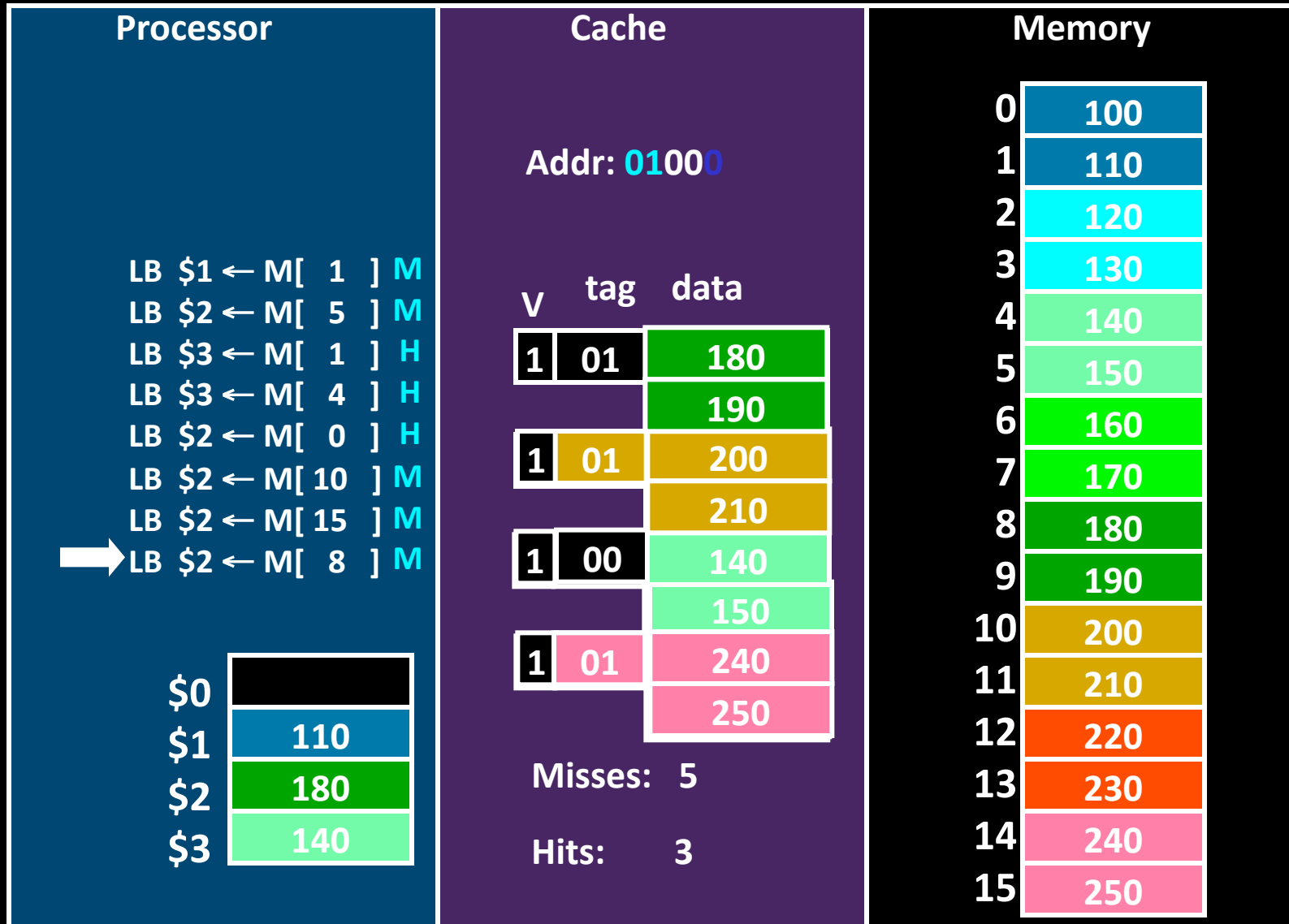
7th Access



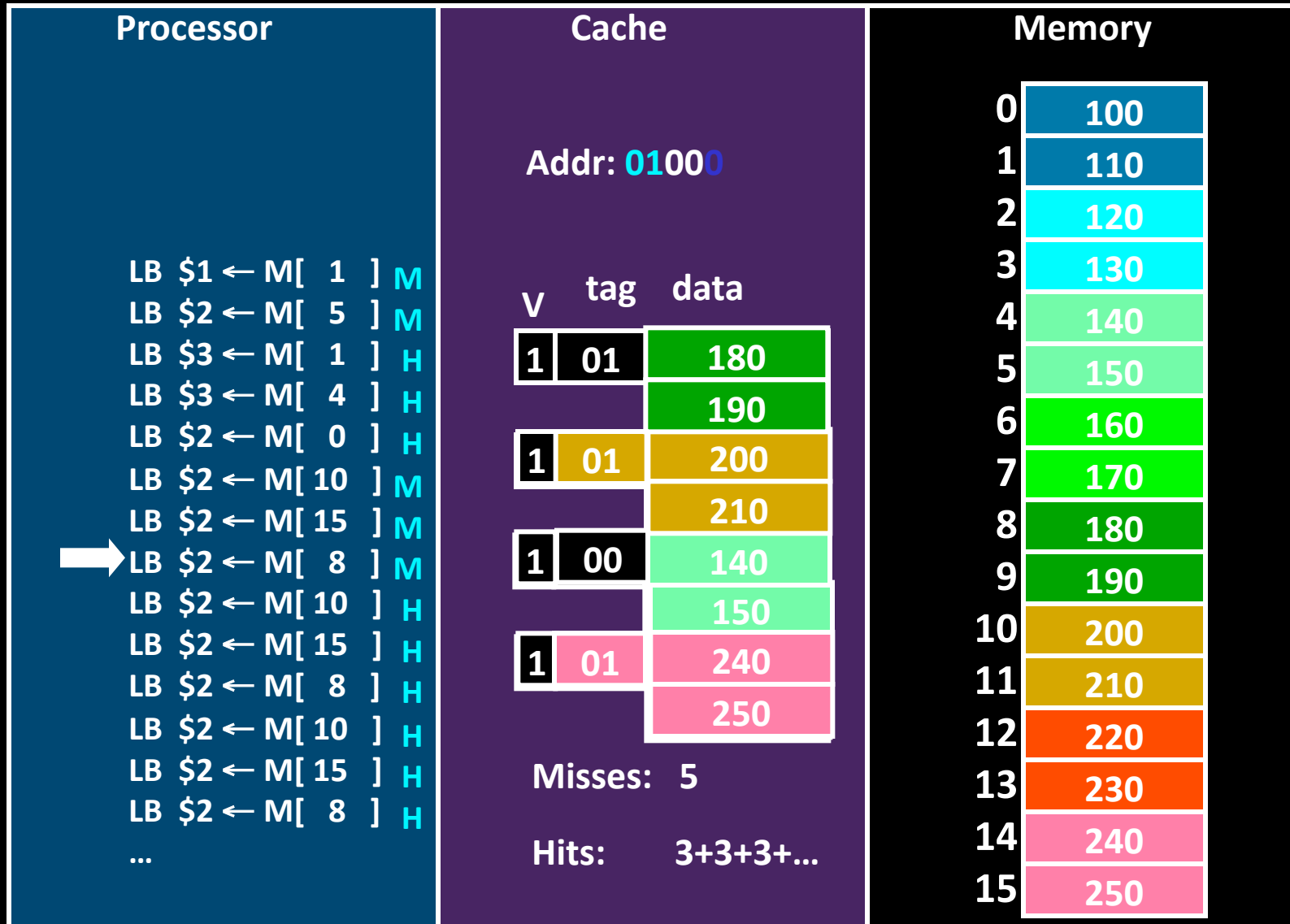
8th Access



8th Access



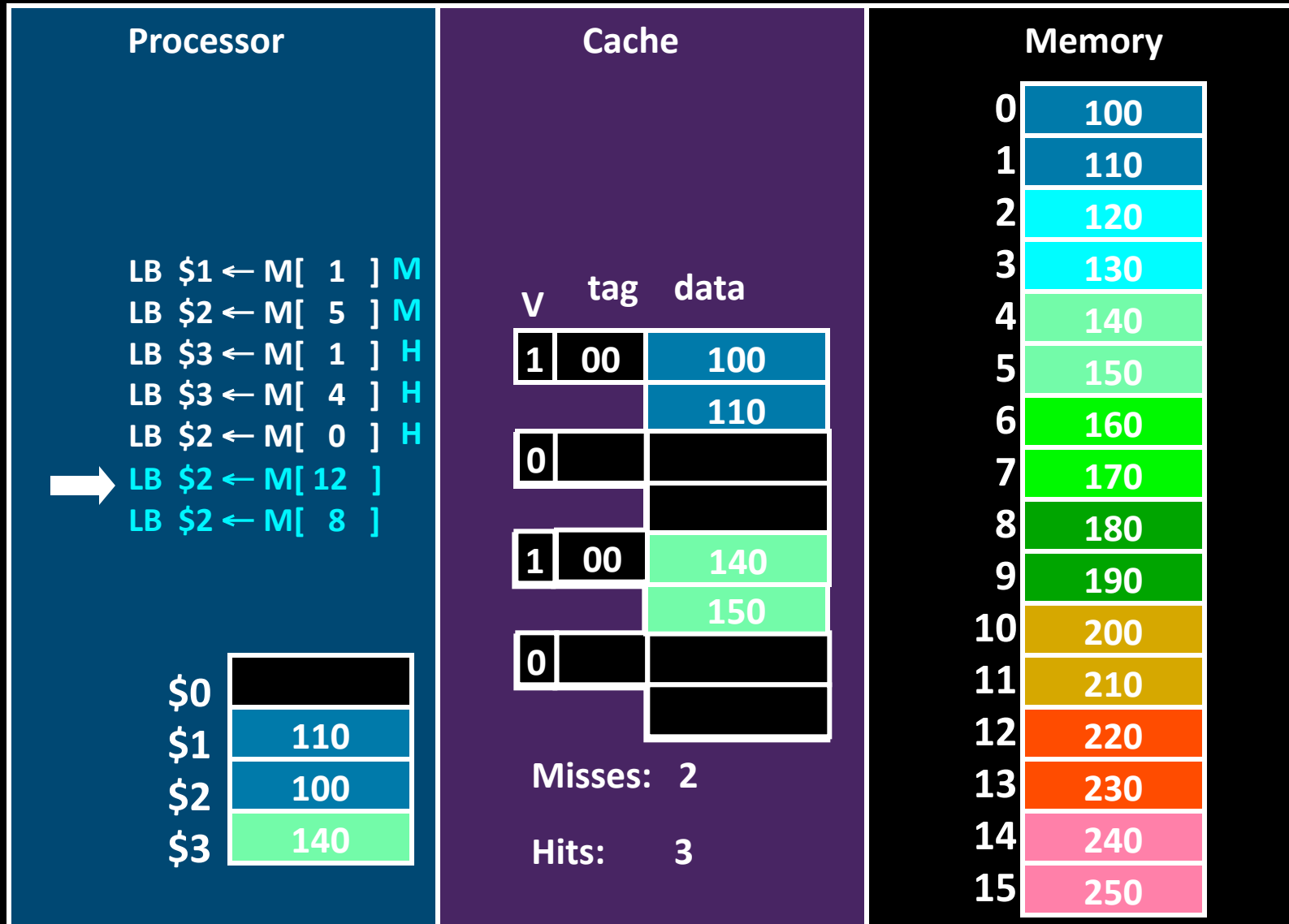
8th Access



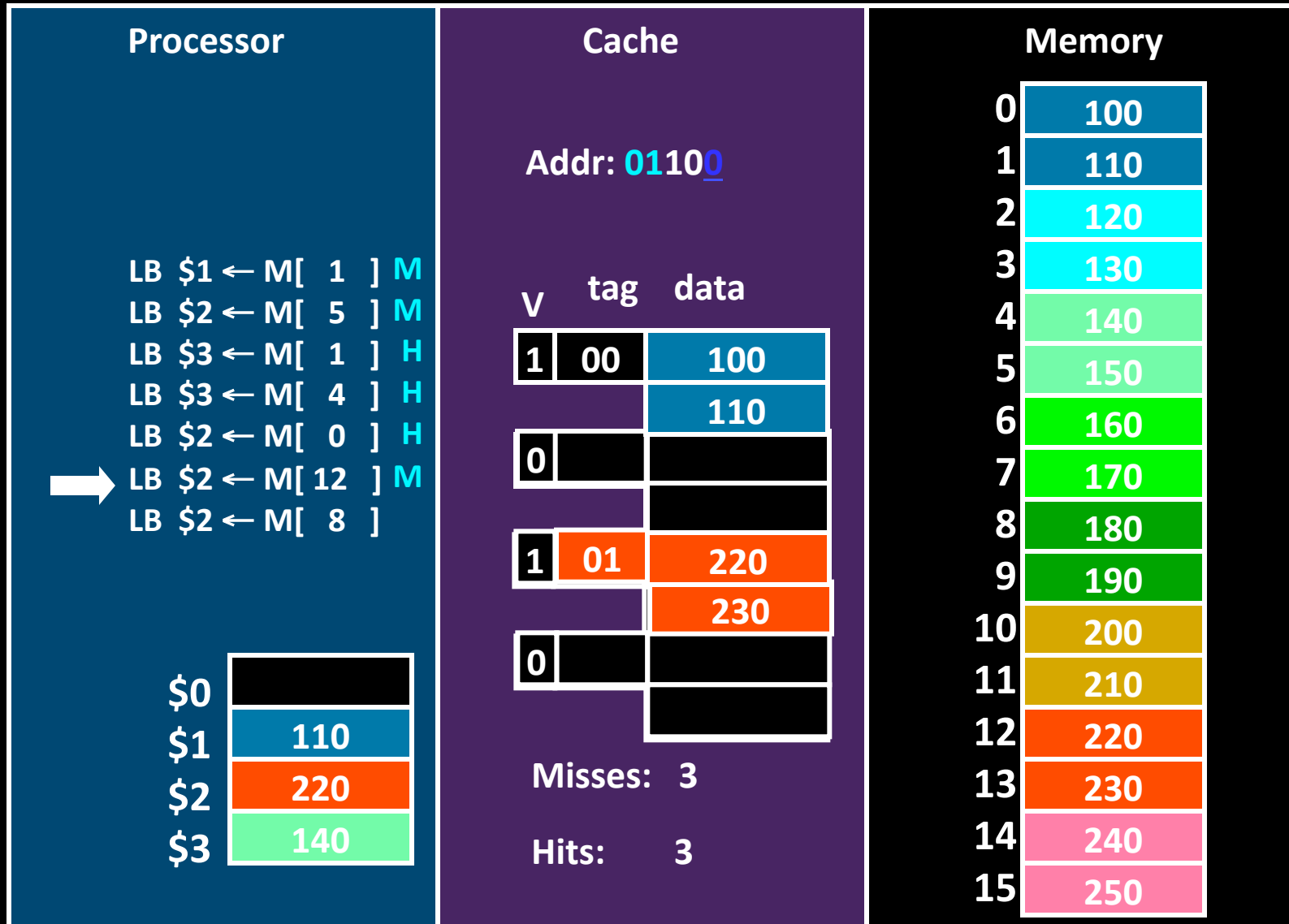
Pathological Case

Direct Mapped Example: 6th Access

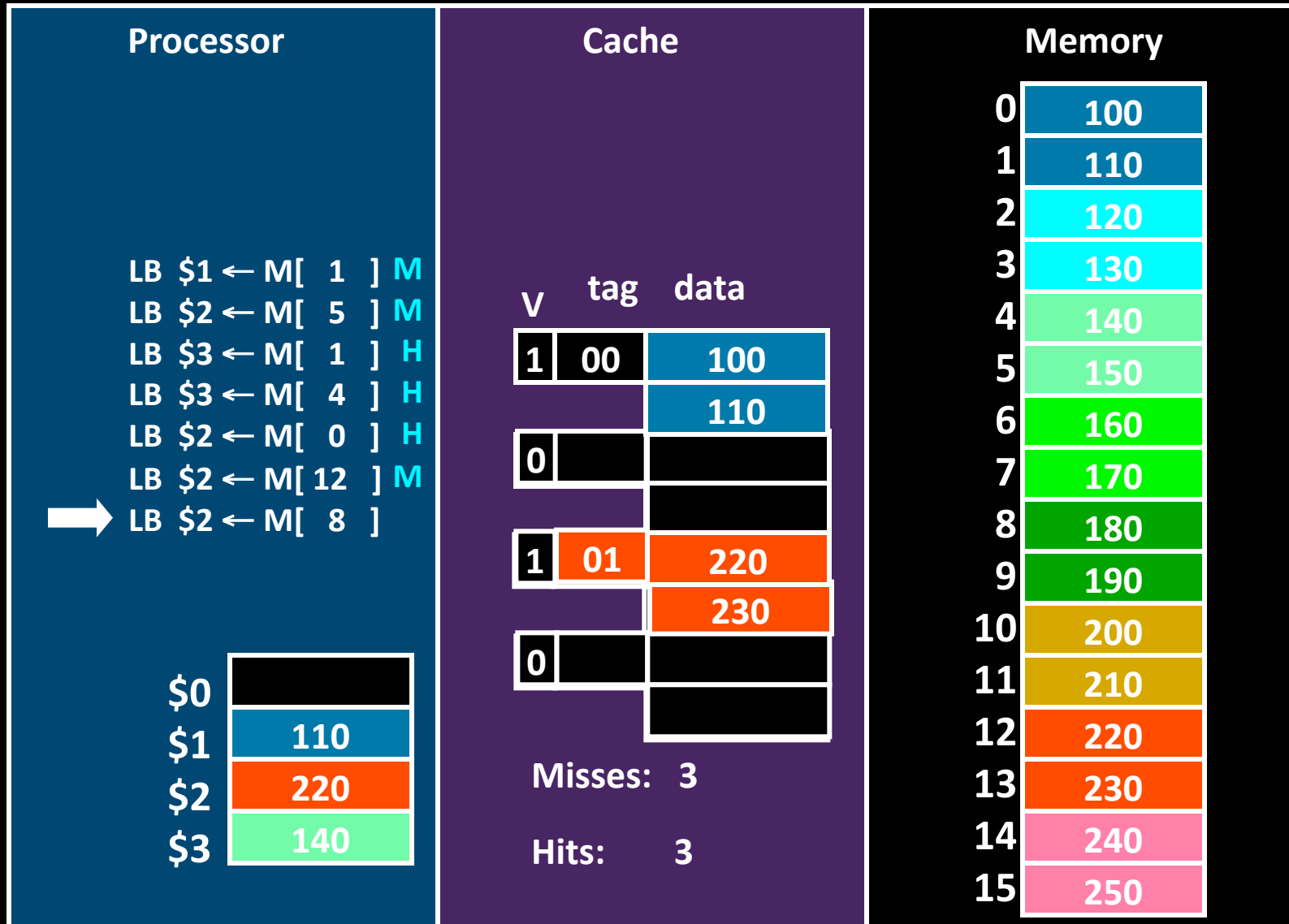
Pathological example



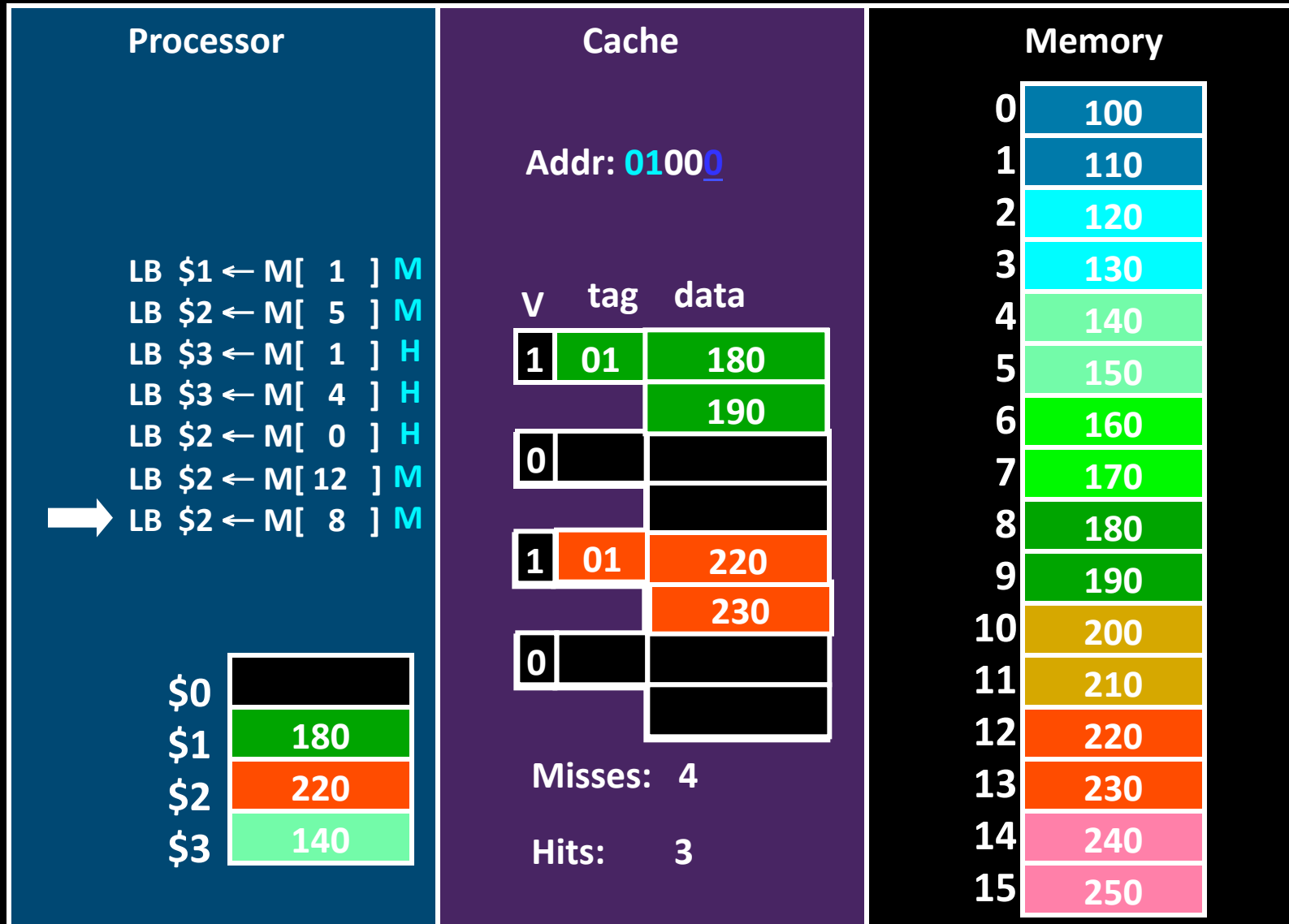
6th Access



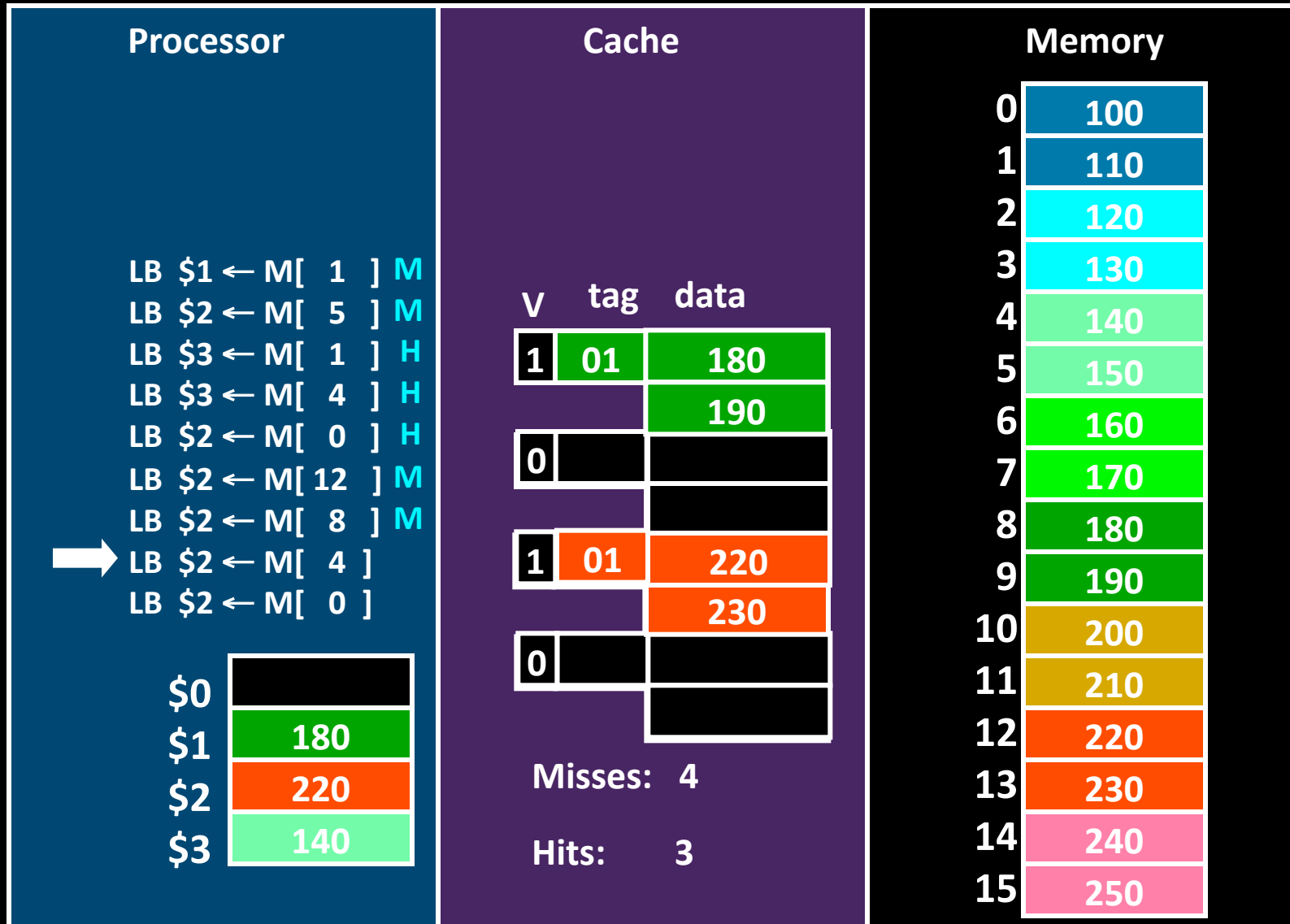
7th Access



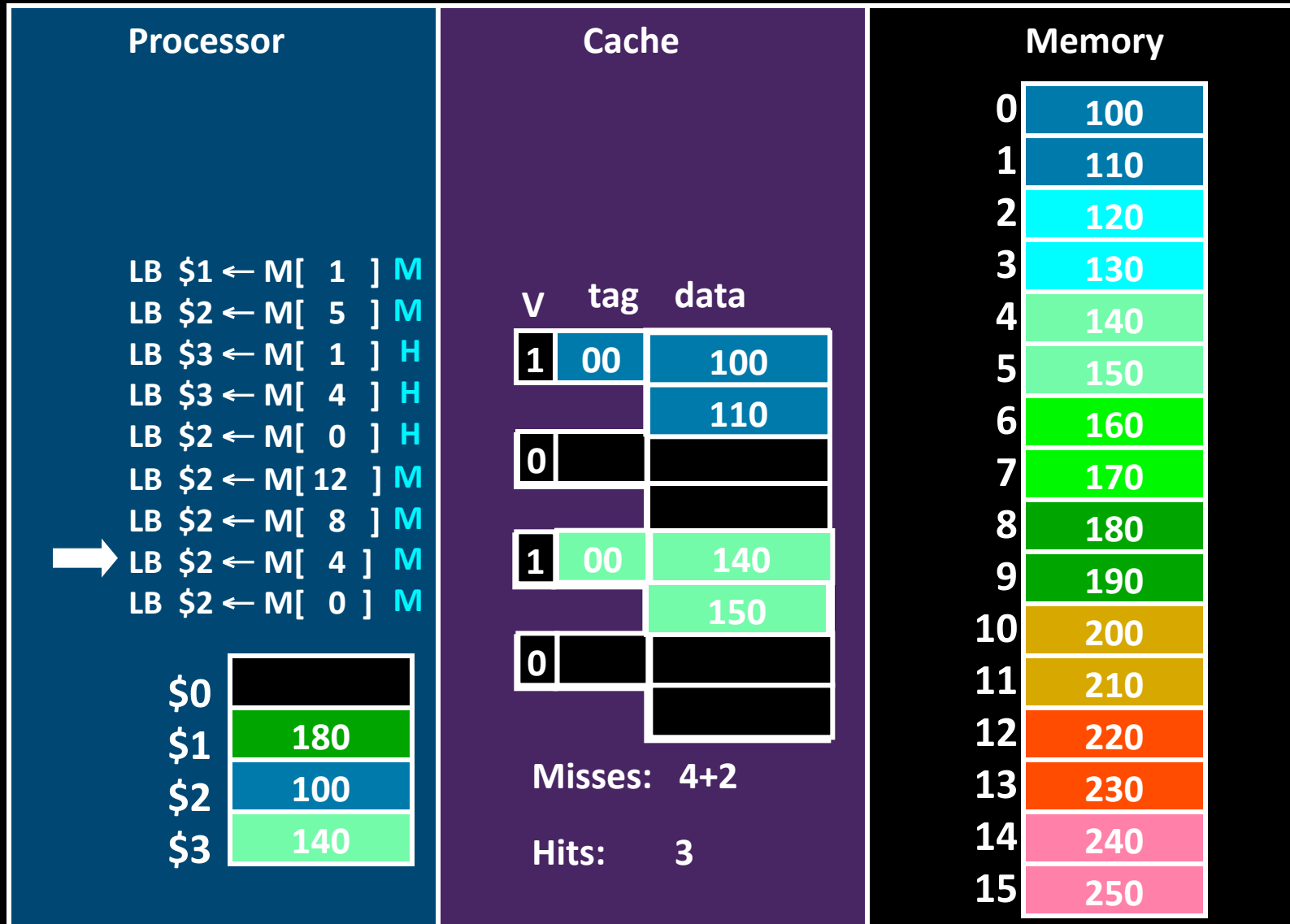
7th Access



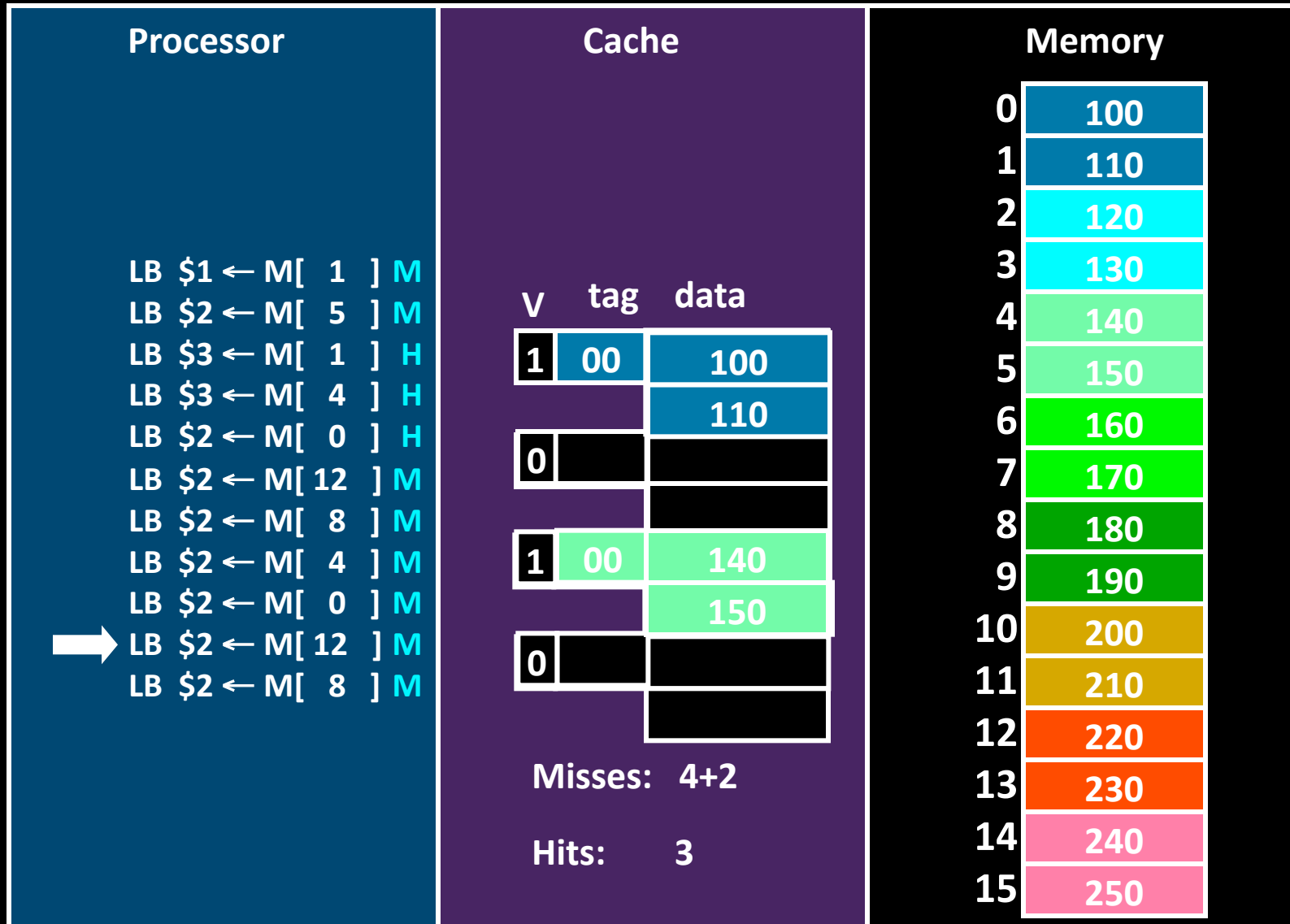
8th and 9th Access



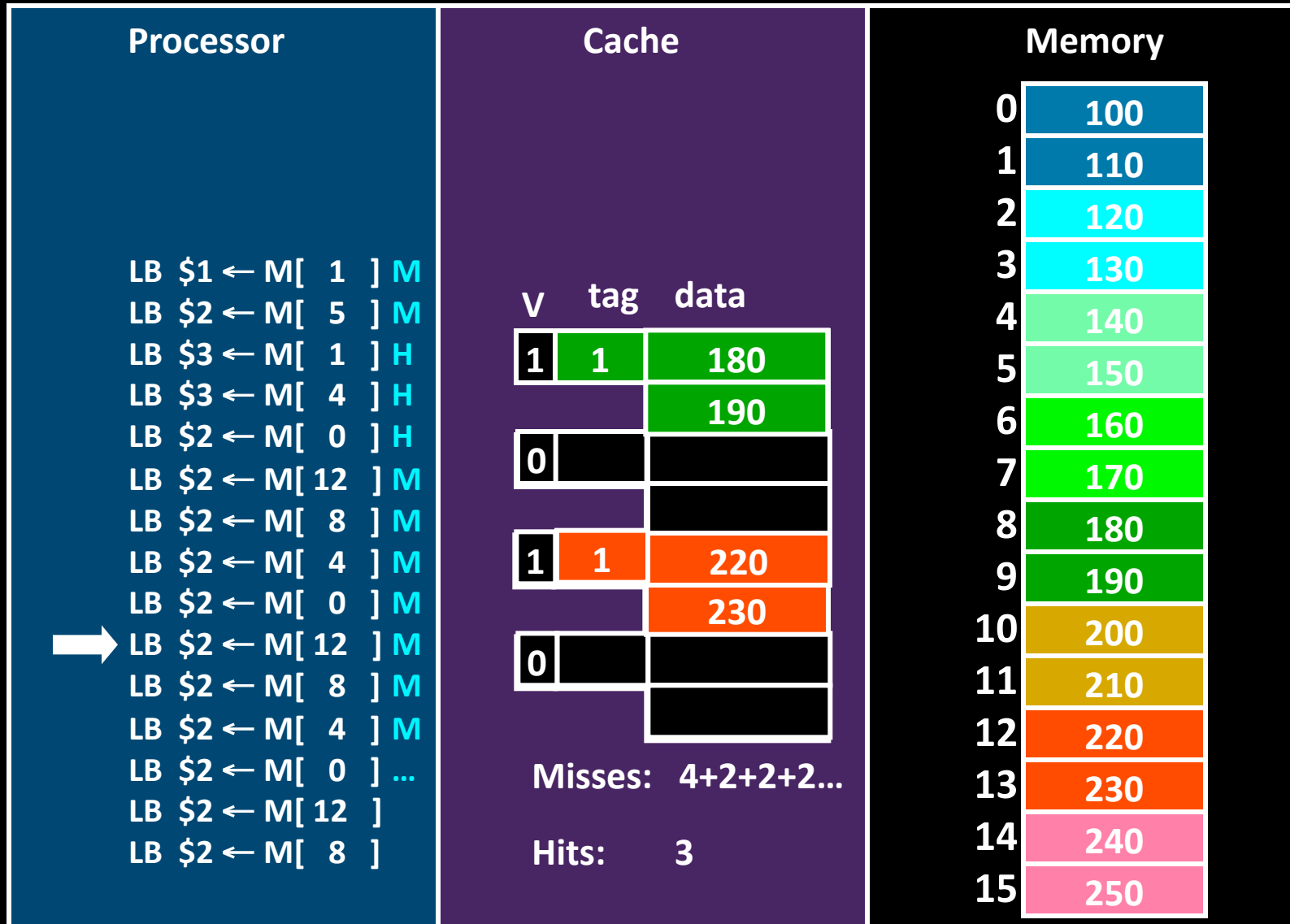
8th and 9th Access



10th and 11th Access



10th and 11th Access



Problem

Working set is not too big for cache

Yet, we can't make it work