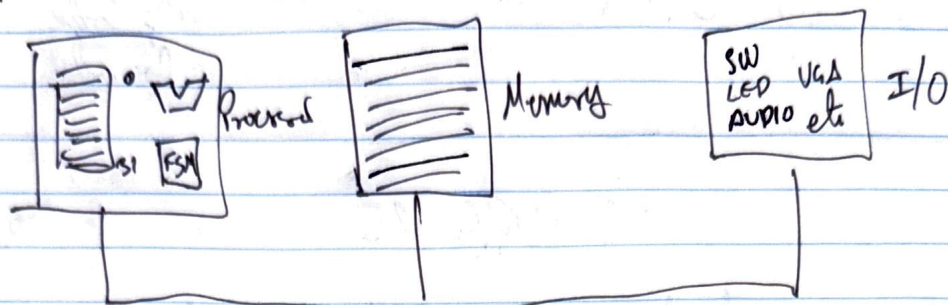


LECTURE 24

06/03

Computer Architecture So Far,...

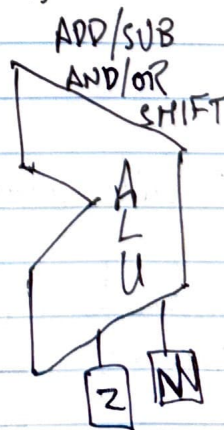


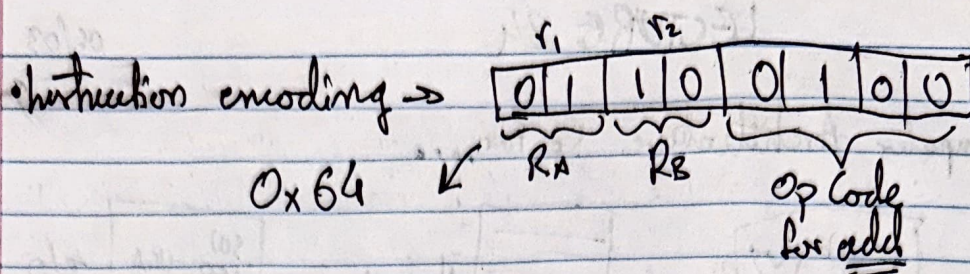
ECE243 Simple Processor

- Four registers $r_0 \rightarrow r_3$, word size 8 bits
- 256 bytes of memory, has a program counter
- Only 10 instructions $\rightarrow$ each takes two arguments
- $NEW \rightarrow Z\text{-bits} = 1$ if result of most recent computation is zero
 - $\hookrightarrow N\text{-bit} = 1$ if result of most recent computation is negative
 - $\hookrightarrow$ Used in conditional branch instructions

- Consider add instruction $\rightarrow$ add $r_1, r_2 \quad r_1 \leftarrow r_1 + r_2$

- 1) Fetch encoded instruction from mem
- 2) Figure out what instruction is
- 3) Compute $R_1 \leftarrow R_1 + R_2$
- 4) Set $N \& Z$
- 5) Increment PC

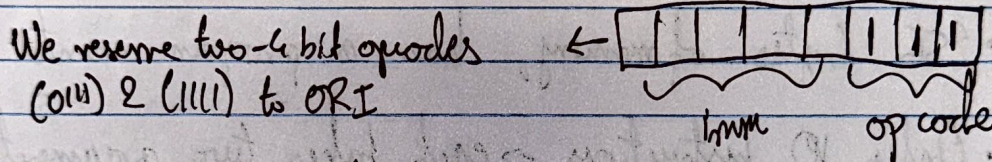




LECTURE 25

10/03

- OR-immediate $\rightarrow$ immediate is 5 bits $\rightarrow$ no room to specify a register
- $\hookrightarrow$ So we always store into R1 and zero-extend the immediate



Conditional instructions

- In Simulix we use result of prev operation to compare to zero
- $\hookrightarrow$ bnz imm4 $\rightarrow$ execution $\rightarrow$ if $(z == \text{imm}4)$ $PC += SE(\text{imm}4)$
else $PC += 1$

- Relative branching gives us movable code

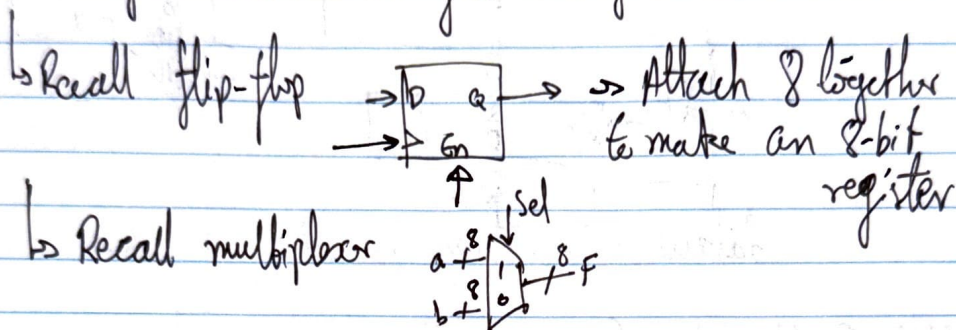
LECTURE 26

12/03

- LGT's write a small assembly program. (we are assembling)
- $\hookrightarrow$ subtracts 2 numbers and branches back to the beginning

Address	Instruction	Encoding	Hex
0 loop:	load R2, (R0) $\#R2 \leftarrow \text{mem}(R0)$	10030000	0x30
1	load R3, (R1) $\#R3 \leftarrow \text{mem}(R1)$	11010000	0xA0
2	sub R3, R3 $\#R2 \leftarrow R2 - R3$	10110110	0xB6
3	brg loop	11011001	0xD9

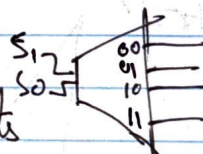
• Recall from ECE241 a few things



LECTURE 27

13/03

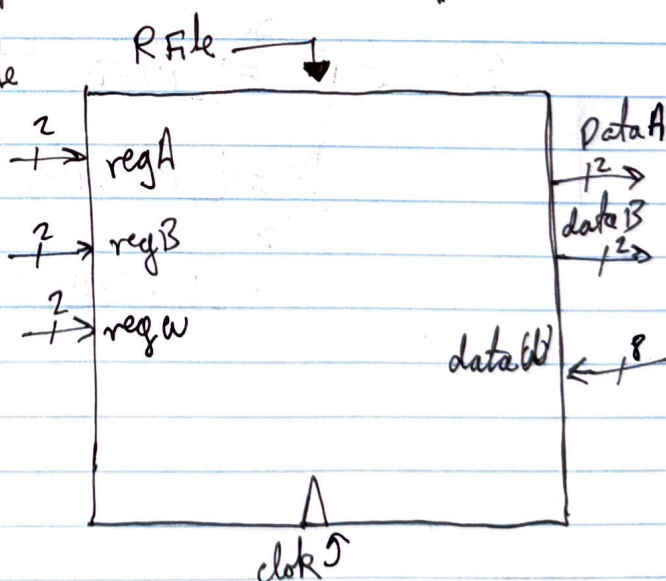
• Recall also from 241 the decoder →

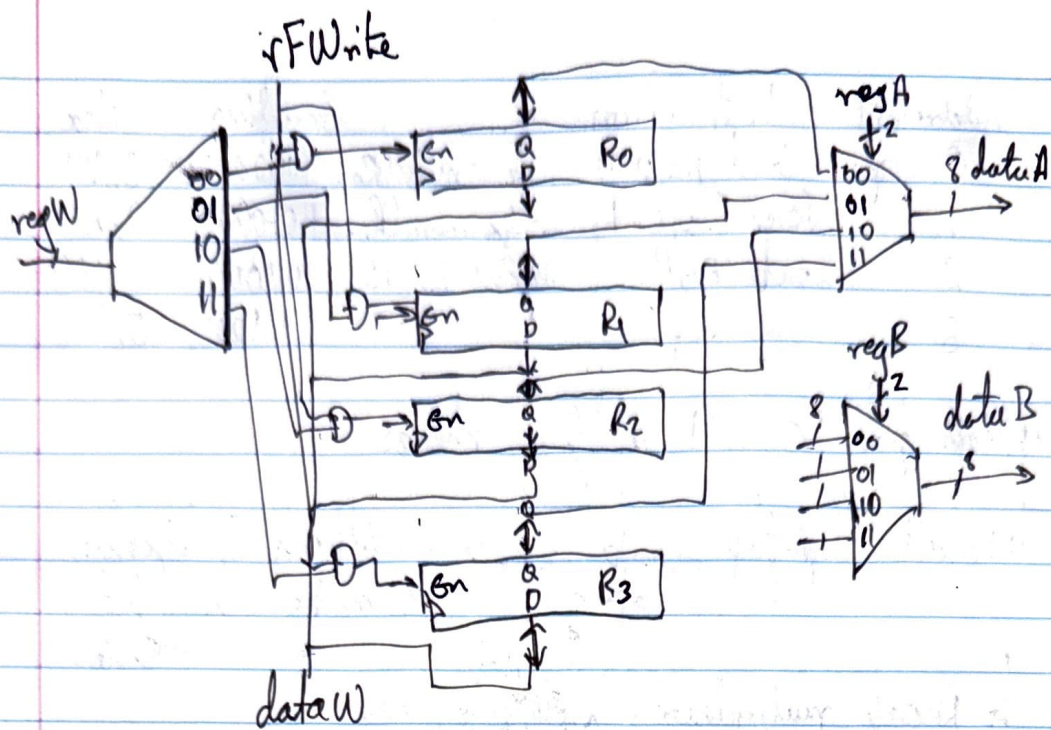


• Building more detail for SimProc elements

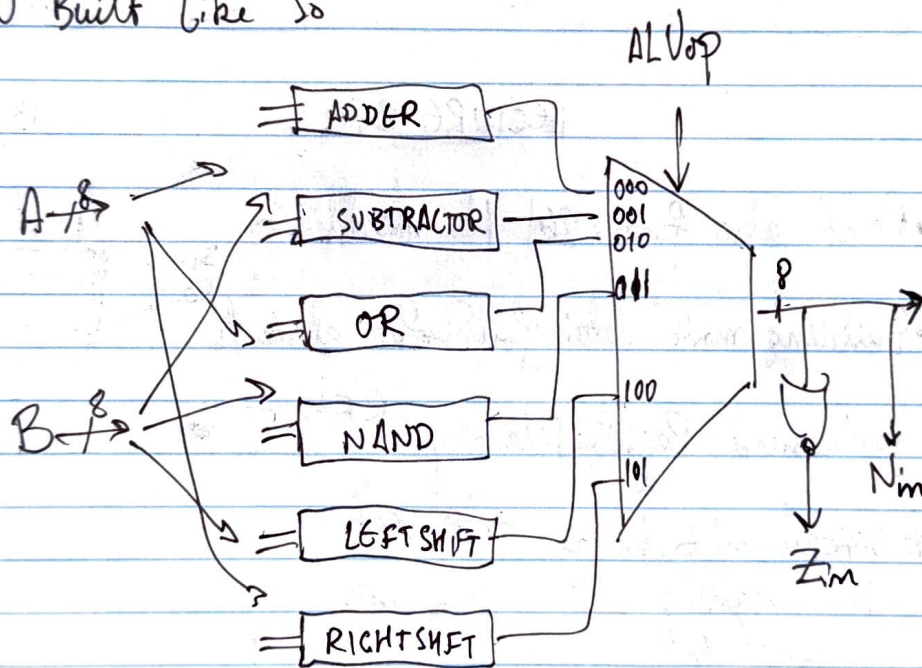
↳ We need Register File

* Careful not to change the registers accidentally

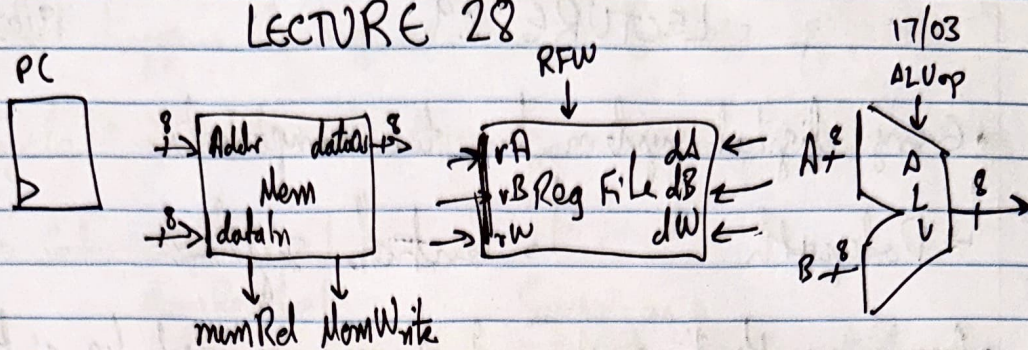




ALU Built Like So



LECTURE 28



• Here's what happens in full execution

↳ ① Fetch op & opcode from memory

↳ Put this into the Instruction Register

↳ ② Decode the instruction → Stored & decoded in IR

↳ ③ Read RA & RB from register

↳ ④ Compute ALUout $\leftarrow RA + RB$

↳ ⑤ Write ALUout back into RA

↳ ⑥ Set Zm & Nm based on the operation's result

↳ ⑦ Compute $PC = PC + 1$

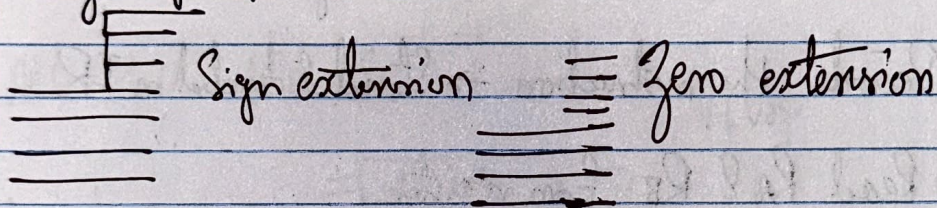
★ Special case happening w/ ORI instruction

↳ Look at last 3-5 mins of lecture to review

LECTURE 29

19/03

- Every digital system has two components
 - ↳ Datapath
 - ↳ Control system
- Anything that comes out of memory must be sorted into either the Instruction Register or Memory Data Register
- Dereferencing data/storing into memory done by interpreting w/ Address & the memsel mux.



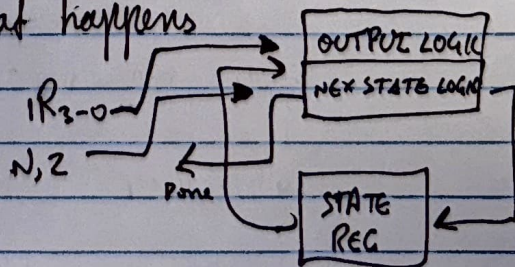
- How does IR know if opcode is 3 or 4 bits?
 - ↳ SimProc automatically responds to X011 as a shift regardless of what X is

LECTURE 30

20/03

Finite State Control of SimProc

- FSM controls which signals are activated & when (on which clock signal) that happens



• What happens during each cycle for ADD/SUB/NAND

① Cycle 1 → Fetch instruction from memory: $IR \leftarrow Mem(PC)$

↳ Set AddrSel = 1

MemRead = 1

IRLoad = 1

Saves us a

↳ cycle of compute!

↳ We can also get $PC \leftarrow PC + 1$ in the same cycle

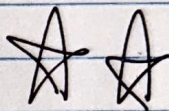
↳ Set ALU_A = 0 ($PC \rightarrow ALU$)

ALU_B = 001 ($ALU_B \rightarrow 1$)

ALU_op = 000 (add)

PCWrite = 1

} Clock-to-Q time
ensures we don't
screw up here
↳ setup & hold time



↓
Shows up in many
interviews

② Cycle 2 → Decode the signal in IR (needs 1 full cycle)

↳ We guess we'll probably need to read R_a/R_b

↳ Set R_asel = 0 (selects R_a to regA)

ABLD = 1

③ Cycle 3 → Add/Sub/NAND specific

↳ Set ALU_A = 1 (select A to ALU)

Set ALU_B = 000 (select B to ALU)

Set ALU_op = one of add/sub/nand

ALU_outLD = 1

FlagWrite = 1

① Cycle 4 $\rightarrow$ Get ALU.out to RA

$\hookrightarrow$ Set RegIn = 0 (connect ALU.out $\rightarrow$ dataIn)
RFWrite = 1

Done = 1 $\rightarrow$ This is for next-state logic

Now let's do "Store RA, (RB)"

$\hookrightarrow$ Only need cycle 3 onwards

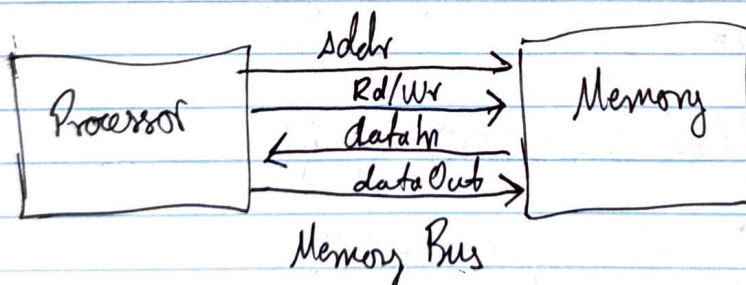
③ Cycle 3 $\rightarrow$ AddrSel = 0
MDRLoad = 1
MemRead = 1

④ Cycle 4 $\rightarrow$ RegIn = 1 (MDRLoad $\rightarrow$ dataIn)
RFWrite = 1
Done = 1

LECTURE 31

24/03

• Introduction to Cache Memory



Processor has two memory access types

$\hookrightarrow$ Read (memory or instructions) $\rightarrow$ both look the same to memory

$\hookrightarrow$ Write (store)

- Some physics for context:

- ↳ Larger memories are slower $\leftrightarrow$ smaller memories faster

- ↳ Longer wires have more resistance & capacitance, making signal propagation slower

- ↳ Decoder time will get slower in larger memories

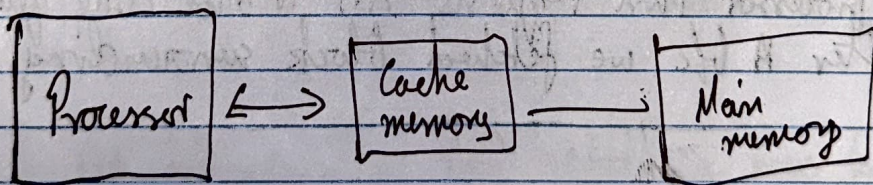
- Interesting point is that as chips get physically bigger they don't necessarily get slower

- 16 GB memory system takes ~300 cycles to respond to memory-access request

How to solve the problem of a fast processor but a big, slow memory?

- ↳ Put a small fast memory in between the processor and main memory & keep the much smaller active memory in between.

- ↳ Called the cache memory



- Sets of instructions and certain data structures benefit from having close addresses

- ↳ Locality of reference:

- ↳ Some sets of instructions called multiple times

• Two types of Locality

↳ Spatial locality $\rightarrow$ physically closer together

↳ Temporal locality $\rightarrow$ Addresses same over time

Operation of a Memory Read Using a Cache

① Processor requests to read data from address A.

② Cache (somehow) looks inside to see if it has a copy of data at address A.

↳ If no, we call a cache MISS. Cache requests data from A (and some chunk of blocks).

↳ Memory responds slowly, cache copies to itself and sends to the processor

↳ If yes, then we call it a cache HIT. Quickly send that data to the processor in just a few cycles

③ If processor then requests A+1 it will come quickly after A b/c we fetched block surrounding A.

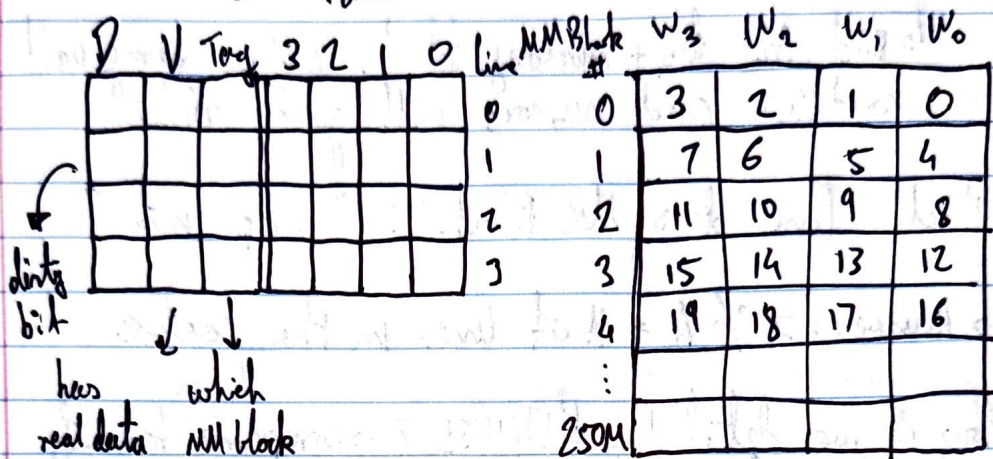
LECTURE 32

26/03

Logistics of the Cache Memory

① Because cache asks for >1 mem location when MISS, we consider main memory as organised into blocks (say 6 words for eg).

- We assume every word has an address



The Direct-Mapped Cache

- Cache has rows & columns like MM, each row called cache line

↳ line size equals main memory block size b/c we take whole blocks into cache lines

- Tag is group of bits that tell you which MM block is in the cache line
- Valid bit tells us if there's real data in the cache or not
- Dirty bit tells us if we write something new to cache, needs to send information back to MM.
- Now consider when an address is requested by the processor → where can the data go?
- ↳ 1 line → "Direct-Mapped" Cache

↳ good b/c we only have to look at 1 tag/line in cache for HIT/MISS

↳ Bad b/c two + memory blocks using same will conflict and memory will get evicted

Which line does MM Block "J" map into?

↳ Answer: $J \% 4 \rightarrow \# \text{ of lines in the cache}$

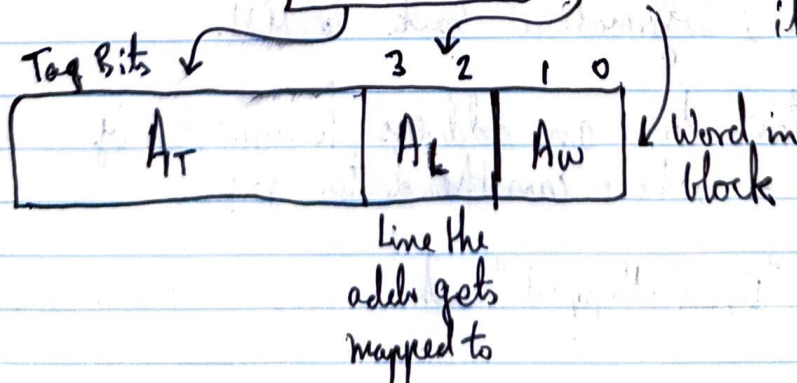
How do we detect the HIT/MISS? $\rightarrow$ compare tag of cache line (if valid) w/ the tag of requested memory block.

Consider this possible memory access request addresses

Block #	Address Bits
	31 30 5 4 3 2 1 0
0	0 0 0 0 0 xx
1	0 0 0 0 1 xx
2	0 0 0 1 0 xx
3	0 0 0 1 1 xx
4	0 0 0 1 0 0 xx
5	0 0 0 1 0 1 xx
6	0 0 0 1 1 0 xx
7	0 0 0 1 1 1 xx
8	0 0 1 0 0 0 xx

Tag (AL)

↓
If matches and is valid then it's a HIT



LECTURE 33

27/03

Consider what happens when the processor makes a read request at MM addr A

↳ cache looks inside itself to see if $\text{tag}(A_L) = A_T$

↳ If yes, and $V(A_L) = 1$, then cache has a HIT

↳ Then the cache sends that word quickly to the processor

↳ otherwise it's a MISS

↳ The line in A_L is valid, it must be evicted
the direct-mapped cache only allows an MM block in one place

↳ The new block is slowly retrieved from MM and copied into cache line A_L

↳ Set $\text{tag}(A_L) = A_T$, $V(A_L) = 1$, $D(A_L) = 1$

Eviction of a cache line

• Method 1 → Write-Back Cache

↳ If there was a write/store to any word in the line, the Dirty Bit will be set

↳ When $D(A_L) = 1$, upon eviction, the line must be written-back so MM is correct

★ ↳ Multiple writes to a line results in only 1 write to MM

Method 2 $\rightarrow$ Write-Through cache

- $\hookrightarrow$ Has no Dirty bit
 - $\hookrightarrow$ All writes to a line are also automatically sent to MM
 - $\hookrightarrow$ Generally undesirable as we may have multiple writes to MM in very short span.
-

LECTURE 34

31/03

New Cache Type: Associative Caches

• Given MM block can map into every cache line

$\hookrightarrow$ Good: Solves the conflict problem in best way possible

$\hookrightarrow$ Bad: Needs a lot of digital circuitry (comparators) to simultaneously check all tags of all lines

• There is a compromise b/w the 2 (direct-mapped & fully associative).

$\hookrightarrow$ Allow each MM block to placed in a smaller set of lines

$\hookrightarrow$ Say "M" different lines $\rightarrow$ called M-Way Set-Associative Cache

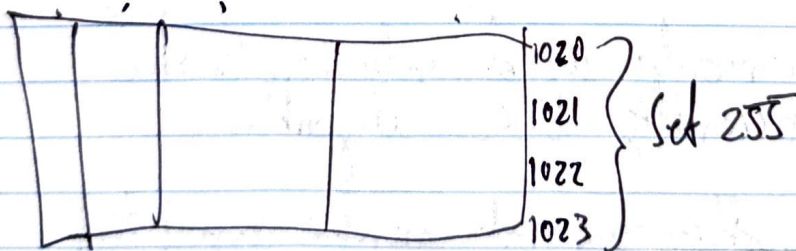
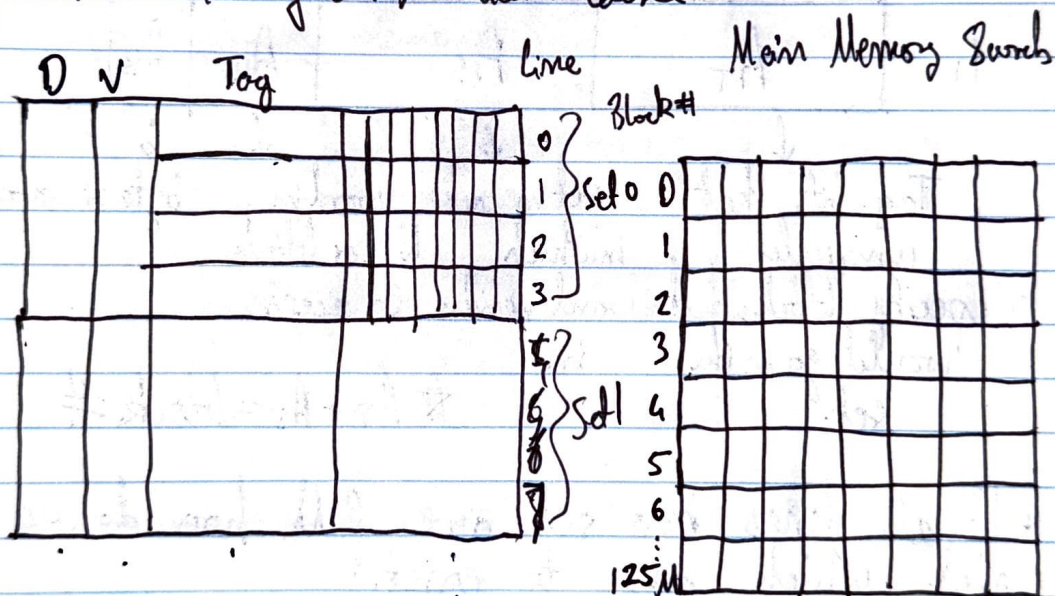
Eg. 1 → A 2-Way Set-Associative Cache allows an MM block to be placed in 2 different cache lines

↳ Those 2 lines are called a set.

Eg. 2 → For NIOS-V (addressability w/ bytes), cache size is 32 Kbytes

↳ Cache has 32 bytes per line, 8 words per line. $\star 32K/64K$ always means nearest power of 2 $\star$

↳ It is a 4-Way Set-Associative Cache



$\star$ For an N-Way Set-Associative Cache, each memory block can be placed into N-different cache lines

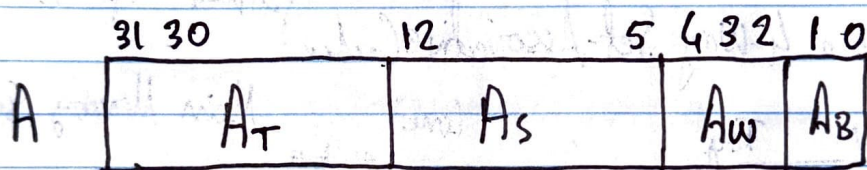
To determine the set-associative cache mapping / HIT/MISS, we need to compute:

$$\hookrightarrow \# \text{ of lines} = \frac{\# \text{ bytes in cache}}{\# \text{ bytes/line}} = \frac{32K}{32} = 1K = \underline{1024}$$

$$\hookrightarrow \# \text{ of sets} = \frac{\# \text{ lines}}{N} = \frac{1024}{4} = 256$$

$\hookrightarrow$ Need $\log_2(256) = 8$ bits to select a set

$\hookrightarrow$ 8 words/line $\rightarrow \log_2(8) = 3$ bits to determine which word in a line/block to address



$\downarrow$ Tag bits that uniquely specify blocks that would go in the set.
 $\downarrow$ set in cache in which word resides in
 $\downarrow$ word in line or block to access
 $\downarrow$ byte or word

$\star A_T + A_S = \text{block \#}$

- If all lines in a set are full, how do we pick which one to evict.

Choosing randomly in general about same as LRU $\swarrow$

$\hookrightarrow$ Most commonly used solution is the Least Recently Used algorithm

$\hookrightarrow$ Have another two bit counter associated w/ each line in the set

$\hookrightarrow$ Set to 0x00 if line is accessed and increment by 1 each time not used

LECTURE 35

02/04

- Spent some time on cache simulator → check notes

Analyzing a Program's Cache Behaviour

- Consider some code examples w/ cache behaviour

Eg 1 → Conflicting access (assume PM w/ BS=8)

```
int A[100], B[100], sum;  
sum = 0;  
for (int i=0; i<100; i++) sum += A[i] + B[i];
```

↳ Suppose every element $A[i]$ maps to same line as $B[i]$

↳ Accessing $B[i]$ after accessing $A[i]$ w/ kick
8 integers out of cache

↳ Happens each iteration of for loop, destroys cache

↳ Soln: maybe split into two loops

Eg 2 → Accessing 2D arrays e.g. frame buffer from lab 7

```
void bad_clear(void) {  
    for (int c=0; c<220; c++) {  
        for (int r=0; r<240; r++) {  
            buf[r][c] = 0x0;  
        }  
    }  
}
```

↳ Horrendous because of row-major array interferes
w/ column-major clear function

Eg 3 $\rightarrow$ Accounting for cache size

```
int A[N], sum, prod;  
for (int i=0; i<N; i++) sum = sum + A[i];  
for (int i=0; i<N; i++) prod = prod * A[i];
```

$\hookrightarrow$ If $N \gg$ cache size, we have to clear cache and redo whole cache retrieval for the second loop

$\hookrightarrow$ Soln: do both sum & prod calculations in same for loop

LECTURE 36

03/04

Performance Analysis of Caches Running a Program

• Recall how memory calls work

$\hookrightarrow$ Processor emits a request to access memory

$\hookrightarrow$ If memory address in the cache, then we have HIT and responds quickly, say in " C " cycles

$\hookrightarrow$ These days $C \approx 4$ cycles

$\hookrightarrow$ If cache MISS, takes M cycles for memory to find that address

$\hookrightarrow$ M is called the miss penalty, ≈ 300 cycles

• Define hit rate = num cache HITS / total num accesses

↳ always $0 \leq h < 1$

↳ We can formulate Δ using $h, C, \& M$ compute average number of cycles to complete a memory access request.

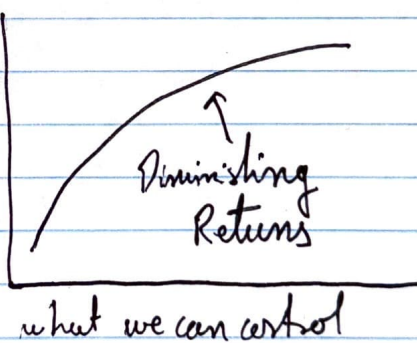
Eg. How many cycles would 1000 accesses take?

$$\frac{1000 \cdot h \cdot C + 1000 \cdot (1-h) \cdot M}{1000}$$

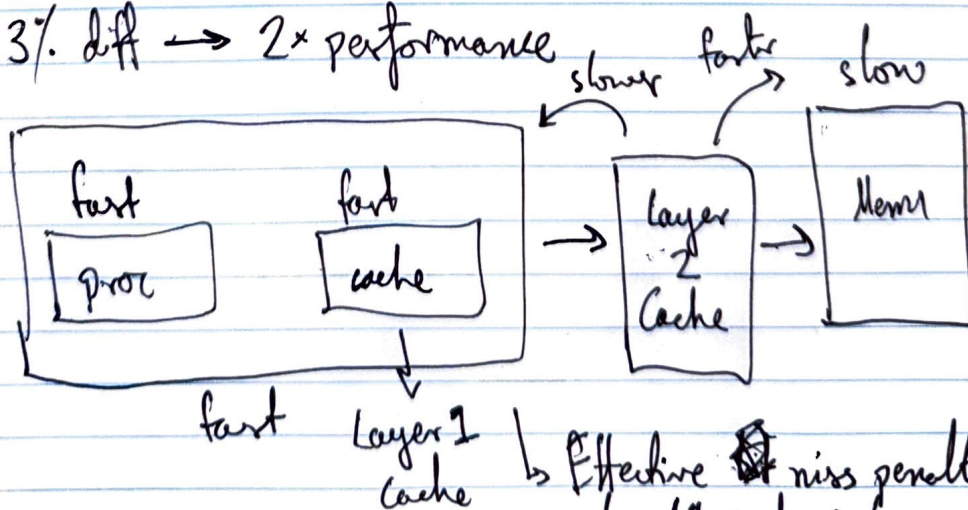
• Have quantity Average Access Cycles (AAC)

$$h \cdot C + (1-h) \cdot M$$

h	AAC	What we want
0.85	48	
0.90	34	
0.95	19	
0.98	9.9	
0.99	7.0	



↳ 3% diff $\rightarrow$ 2x performance



↳ Effective miss penalty seen by L1 cache is lower

• Because instruction memory accesses have different patterns than data accesses, the L1 cache is split into

↳ L1 instruction cache } Architect can design these differently depending on requirements
↳ L1 data cache

• Discourse on Intel's reign and arrival of competition

• Discourse on China-Taiwan-TSMC relations