

Exceptions → **Exception**: unusual internal event caused by prog. during exec → **interrupt**: external event
 → **trap**: forced transfer of control caused by exc. or interrupt & not all excap. cause traps
 → **Asynch interr.**: I/O asserts on prioritized interrupt req lines, choose to process interr, stops at instr I_i , completing until I_{i-1} (**precise interrupt**); & EPC to where it didn't W/commit
 ↳ **Trap handler**: save EPC before enabling interr, ERET (enable + vstart) **Synch trap**: particular instr, usually restart

S-stage: hold exception flags until **commit point** (M stage, speculation lost)
 & only wr at commit, else throw away & flush & launch handler & **Bypassing** (allows uncommitted instr use of results by follow)
 FFDX M (W) - commit & FDX M (M) - NOT commit & FDX
 & raise → wait for all to get to F- then start
Imprecise: FDX M1 M2 W X outputs M1 M2 FDX M1 M2 W
 Reduce latency: unvail, kill precisely

Latency: t taken for single op to start to finish **Bandwidth**: rate at which ops can be performed
Occupancy: t during which unit is blocked on op (struct)

Memory **DRAM** (optimized for density) **SRAM** (optimized for speed), not destructive
 CPU ↔ SRAM ↔ DRAM & Temporal: for self, ref. again & Spatial: near locations
 3 predictable forms

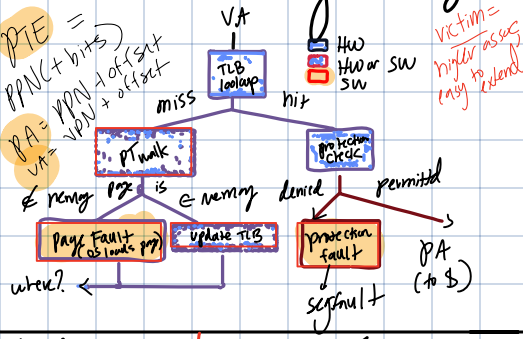
\$ **Replacement** & **Random** (Fast) & **LRU** (complex) & **FIFO** (high assoc) & **NMRV** (FIFO except most recent)
 \$ **Policies** (cache hit & write-thru (\$+mem) & write back (\$ only, eviction spawns wr) } wrthru + wr, no
 Cache miss & No-write alloc - wr to mem & write alloc - fetch into \$ (only to mem if evicted) } wr back + wr alloc

C's: **larger line size**: ↓ compulsory ↑ conflict **larger \$ size**: ↓ capacity/conflict ↑ hit **higher assoc**: ↓ conflict ↑ hit

Write buffer: hold wr before \$, reduce read MP & reduce wr penalties (not write) & reduce read penalties (not write) Misses private
 & **local miss** = miss in \$ / access to \$ **global miss** = miss in \$ / CPU mem access (MPI) = #misses / #instr

Prefetching **Accuracy** (useful / prefetches) **coverage** (useful / total unp. access) **Time** (on-time / total pref.)
 Algos: **Next-line** (next N \$ lines after comp. miss) **strided** (after saving N w/ dist D → + C + D) bit 2N

Virtual Memory **Page**: fixed size, internal frag, efficient & **Segment**: variable, ext.



ex 4 GiB VA, 4 KiB pages, 4 byte PTEs
 ↳ $2^{VPN} \times \text{page size} = \text{PA}$
 ↳ **page offset** = $\log_2(\text{page size in bytes}) = \log_2(4096) = 12$ size address-able
 ↳ **VPN bits** = **Size of Mem addr** - **page offset bits** = $32 - 12 = 20$
 ↳ **# virtual pages** = $2^{\text{VPN bits}} = 2^{20}$ PTEs & $\text{VPN} = 2^{\text{VA}} / 2^{\text{size}}$
 ↳ 1 page mapped to PA → 1 valid PTE
 ↳ **linear page table** → $2^{20} (\text{PTEs}) \times 4 \text{ B (size)} = 4 \text{ MiB PT}$
 ↳ **2 level**: 1 page mapped to PA → 2 valid PTEs → $2^{10} \text{ PTEs} \times 4$ } 12

VA: 0x58 → OXS = 0101 → 01 = index 1 (64 bit words → 0x00 + 8) → go to 0x000 → apply 12/3 → PA
 Final PA: Content + offset ex 0x17 content → 0x178 final & Bring in pages for faults

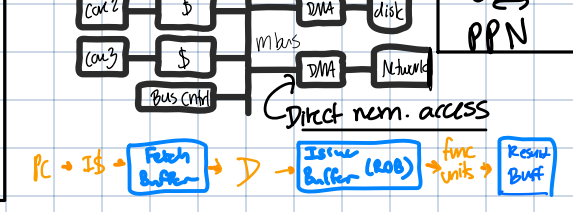
AMAT: HT: incr if crit path ↑ **MP**: increases w/ Block size / outv num, **MR**: decr. w/ bigger tag

ex **Matrix**: Bits 0 to 11 offset + index → 2^{11} VS. + 4096 bytes btw col → 2^{12} → index

TLB: \$ for PTEs, fully assoc & TLB miss ≠ Page Fault & TLB lookup = tag check
 TLB entry: valid/dirty bits, tag, PPN → $\text{PPN} + \text{offset} = \text{PA}$, use VPN for tag

VIVT: cache index & cache tag derived from VA **VIPT**: VA for index PA for tag → parallelism

Digits		2^2	4	2^8	156	Ki	10
0000	0	0101	5	1010	A		
0001	1	0110	6	1011	B		
0010	2	0111	7	1100	C		
0011	3	1000	8	1101	D		
0100	4	1001	9	1110	E		
					F		
		2^3	8	2^9	512	Mi	20
		2^4	16	2^{10}	1024	Gi	30
		2^5	32	2^{11}	2048	Ti	40
		2^6	64	2^{12}	4096	Pi	50
		2^7	128	2^{13}	8192	Ei	60



TLB: Tag (VPN) ↔ PPN
RISC instrs
 lw rd, rsl beq = bge >
 sw rsl, rsl2 bre ≠ slt <
 mv rd, rsl blt < slti < imm
 bge ≥ sltu < v
 bltu < sltiu < imm

Midterm 2 - CS152 - Katrina Sharonin - SID: 3036613337

OoO (ILP) Fetch: instr bits retrieved from $I\$$ → Decode: instr dispatched to issue buff → **Dispatch**: place into ROB, ready to issue → **Issue**: send to exec units → **Completes**: results & exp. available → **commit**: "return" exit ROB @ store modifies state **★ Always commit in order for precision!**

RAW: r_1 op r_2 r_3 op r_4 **WAR**: r_3 op r_2 r_1 op r_2 **WAW**: r_1 op r_2 r_3 op r_2 **★ WAW/WAR in OoO unless in order issue**

VS INO: issue stalls on RAW dependencies / struct / WAR/WAW hazards
OoO dispatched in order to reservation station / instr buffer → wait for ops → dispatch / issue out of order

OoO Design: reservation station (store instr that hasn't issued) → in ROB vs. indep. (issue window) **Expensive to build**

Renaming Design: Tags + data in reserve station VS Tags only in reserv. st., data in unified phys reg

ROB: executing, not committed **Rename Table**: lookup ISA (architectural, x6) ↔ phys (p6)

Register renaming ★ Make WAR/WAW impossible ★ sd, branches not renamed, KO_3

Data-in-ROB = Circular (head = oldest) • p set on result

- Set v on dispatch
- Completion → search src tags → set p & copy data into src on tag match
- Set f on issue
- Commit → check exp. → copy into reg file and 0/1/err → flush, set free oldest, KO_3

★ One rename register per each ROB entry

Unified Phys Rename all arch registers into

Single phys reg file (temporary commit KO_3 , completion)

Rollback: go in reverse order ★ expensive / inefficient since causes a lot of work
superscalar renaming is not possible (2+ instr issued in parallel, RAW → "rear")

Rename	ROB	Free
X0 p3	use ex op p1 PR1 P2 PR2	rd lrd pld
X1 p0	V x ld p p7	x1 p8 p0

★ **Speculative store buff / ld/st queue**: allot in decode → don't know addr & store yet → hold speculative data. Commit when oldest instr & both addr/data available → clear S bit, move data to \$, or

Load Bypass: $SSB \leftarrow \text{pref.}$ Load addr → $L\$$ ★ Use youngest sd that's older than load

VS In order Mem Queue: sd, ld: X in program order, ld/st can't leave ROB until older done X, still OoO spec. KO_3

VS Conservative OoO: sd, ld: don't execute load if any prev. st addr not known, cmp ld addr w/ all uncommitted st

VS Addr spec: sd x1(x2) ld x3(x4): guess x4 = x2, X load before st addr known → hold uncommitted INO → flush if bad

★ **high penalty if bad guess**

★ **aliasing**

★ **IF UK in SSB, override \$**. IF not in SSB, use \$, else VK (VK = UK could impact any)

VS (speculative) assume last known / in \$ even if UK marked

→ **Control Flow Penalty**: reads to jumping back → sol's: (SW) elim br/roll, early cond. eval (HW) delay, speculate ex br. predict

Branch Prediction Directional (dest addr = inspecting instr bits + PC) **VS** Indirectional (reads other reg. ex jr x1)

Conditional (deterministic using input) **VS** Uncond. (return/go to) **VS** Forward (jumping) **VS** Backward (lagging, 90%)

Temporal (the way the same br resches indicates → BHT) **VS** Spatial (several br correlated → Global ctr)

Branch Hist Table (BHT): FSM to predict, update if correct/not; picked using lower bits of PC (or don't alias)

History register: records dir of last N branches (ala global history)

★ **Branch predict VS both dirs** → not efficient → focus rsrcs on predicted dir

★ **Limits**: BHT at D but must wait till X → penalty, if did other instr & row has to flush

→ **Solution**: BTB! BTB in E → inserts addr for prediction, BHT at D + use BTB value

Branch Target Buff (BTB): Entry PC (tag), Valid, Target + predicted PC. IF Tag = match → use, else PC + 4

★ **helps with** cond., direct, indirect ★ only taken branches & jumps held in BTB, next PC calc before br F/D

VS Jump register (JR): (VS BTB) reg designated where program should cont, calcs location, good for:

Switch statements (indir) **VS** BTB (if some case used N+) **VS** BTB (some func called ex virtual func) **VS** BTB (ret's (ret addr) to some place core)

Subroutine Return Stack / Return Addr stack (RAS): store retr. addr as called, ^{accelerate subr. rets} more acc than BTB
 (x) f(a) & f(b); f(b) & f(c); → $\begin{bmatrix} f(c) \\ f(b) \end{bmatrix}$ → ① Lookup BTB for jmp ② push ret. addr ③ ^{return} pop stack ④ jmp back to addr

InO: no instr issued after Br can WB before branch resolves VS OoO can complete OoO but cannot commit

VLIW (V. long instr Word) force compiler to schedule so that HW doesn't worry abt interlocks

VLIW = N + Bundles, each is a type specific slot @ $\approx 2^{mul}$ ^{1 add} ★ Guarantee intra-instruct. parallelism

Loop unrolling: mul iter in one big core, faster stride, a lot of overhead

SW pipelining: same # iter, only 1 time overhead, modify instr to hide latency

★ Write out two iterations: 1st enters, 2nd starts in loop, then last is covered by loop

Best to introduce if MEM in loop: is not fully utilized (spaces exist)

VLIW Issues: obj-code compatibility, wasteful size, unpredictable latency is dangerous to it ^{ex. \$}

Compilation techniques: static schedule (fixed/no prediction), SW pipe, unroll, trace scheduling

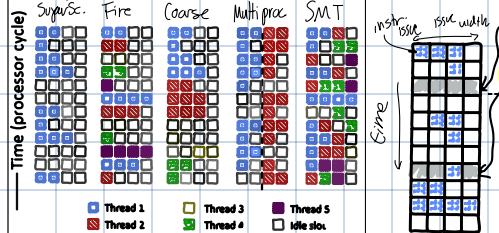
Predicated execution: Mispredicted br limit ILP → eliminate hard-predict br via "if-else encoded instr"

Instr conditionally execute; avoid var. len latency: If cond → do this else → NOP

Trace Scheduling Trace: path of blocks for most taken br path, schedule whole trace at once

Multithreading (TLP) Incr. Thruput (more instr sent) Worsen latency (competing concurrently for res)

★ Doesn't benefit from br. resolution latency (can hide); arch. identical to multicore system to ISA



Vertical waste: no instr issued in single cycle (Reduce w/ OoO via ILP, InO stalls)

Horizontal waste: not all exec units in use in cycle ★ Constrained by # phys reg

Fire-grain: every cycle, ^{hide thruput losses for stalls, don't need low overhead instr, split thr into its own core, reduce hor. limits upper thruput}

Coarse: when blocked, reduce penalty of high cost stalls

Multi-P/CMP: split thr into its own core, reduce hor. limits upper thruput

SMT: smart interleave, full thruput
 ROB size const, storage reg. not changed, low cost + add ^{hard + vent needs multi-issr processor}

i-count policy: for SMT, fetch thread that has least instr in flight (if stalling/in flight, will get stuck again...)

Sharing Resources vs. Dup: (Duplicated) PC, Rename, ^{Arch. File reg.} (Shared) Fetch unit, phys Regfile, ^{issr window} Func units, ROB

```
loop: ld $s0, A($1)
      ld $s1, B($2)
      addi $s2, $s0, 1
      addi $s3, $s1, 1
      fmul $t0, $s0, $s1
      fadd $s4, $t0, $s2
      sd $s4, S($3)
      bnez $s0, loop
```

Fire grain
 ① Find max
 ② $(S_{max}) - (N_{min}) \geq 0$
 ③ Find N

Data-Dep
 ① Ellipse
 ② Steady state max jump
 ③ $\text{ceil}(\frac{S_0 - 4}{7}) + 1$

Vector ISA
 $\text{vsetvli } t0, a0, e8, m2, ta, ma \text{ (Config vlen \& rtype)}$
 $\text{vlsr } v1, v2, rs1, rs2, vm \text{ (32 bit shiftd load)}$
 $\text{vmsl } v0, v2, x0 \text{ (set if less than (signed), } v0.t = \text{max})}$
 $\text{vlvxei } 32, v2, x0, v1 \text{ (load a[idx(i)] use element of v1 to index into a)}$
 $\text{vmmul } v1, v2, v2, a4, v0.t \text{ (Scale * a[idx(i)] IF cond. true)}$
 $\text{vssr } 32, v2, v2, a2, v0.t \text{ (Store into b[dilation * i])}$

Vectors (DLP)
 Compact, Scalable, expressive
 Packed SIMD: fixed vector len, limited instr set, Scatter/gather, align
 Key: sup. dispatch to stay busy; loop unrolling ↑ reg. pressure

for (unsigned int i=0; i<N; i++)
 if (a[idx(i)] < 0) {
 b[dilation * i] = scale * a[idx(i)];

Vectors Cont. Vector len register = # valid ^{up to} MAX VL / VLMAX
 Standard elem width (SEW): default width of each vec elem
 Len Multiplier (LMUL): int power of 2 grouping ^{16 scalar vectors @ = 2} vlen & vlen+1
 VL = # lanes → 32 elem vector, 8 lanes → 4 vectors issued first

```
loop: vsetvli t0, a0, e32, m1, ta, mu
      vle32v v1, a3
      vll.vi v1, v1, 2
      vluxe32v v2, a1, v1
      vmsl.vx v0, v2, x0
      vmmul.vx v2, v2, a4, v0.t
      vssr 32, v2, v2, a2, v0.t
      sub a0, a0, #4
      mul t1, t0, a5
      slli t0, t0, 2
      add a3, a3, #4
      add a2, a2, #1
      bnez a0, loop
and: ret
```

Mem system: (Banking) split addresses out, go to sep bank w/ high probab. ^{not dup} ^{dup} ^{no chain} ^{chain}
 Chaining: Bypass vals early; two Rs but separate func units

Strip-mining: finite Vsize → split across mul vs; Best when iter len > Max VL

Vector reduction: BinTree reduction, recursive aggregation
 $\text{sum} += A[i] \rightarrow \text{sum} = \text{sum} + A[i] \text{ dyphic } A[i-1] + A[i-2] \dots$

STMD (vector) MIMD (thrd) SISR (ooo) GPU: SMT (multi thr) = DLP x TLP

Thrad blocks: range over iteration space Warp: sub blocks not visible to user, micro-ops usually together

VA PA translated addr vs. using VA to address data in \$ → Mul VA + single PA (Aliasing: sim w/ reads)
 ★ if issr → incr. assoc (less index) incr. page size

VIPT: offset bits = virtual index → Ptag & check ★ Page offset == offset + index bits, dec L2 to check

Final Cheat Sheet

\$ coherence

Coherence: legal values a single mem addr can return ^{transient if not atomic}

Write policies: write-back: only \$, mem or flush, write thru: main mem

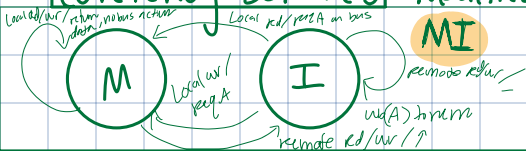
Miss policies: No-write-alloc: only to main mem wr, no \$ wr.

Write-alloc: fetch into \$, avoid wr to DRAM unless evicted

Inclusive / Exclusive: Inclusive ($L1 = L2$), Exclusive ($L1 \neq L2$)

Coherency schemes Modified (I own, not shared) → Owned (shared), exclusive (clean)

MESI to reduce broadcast
if local wr in S state (we might be the only ones)
↑ VS.



causes threads to invalidate each other



Sharing status: Dir based: block of mem (centralized - SymMMP) (distributed: SGI, diff nodes)

Snooping: \$ line, monitor bus: reads totally ad-hoc broadcast network, can go w/o wire bus

Memory access: producer-consumer: wait, use bit to indicate → OoO may give stale data

Mutual exclusion: one process at a time, use lock

row ex coherence issue

ao = a1 = A

preserve program order: Rather W should be up to date

3: W to, (ao) 2: W to, (ao)

1, 2, 3

write serialization: 2 W to same loc. by any 2

1: W to, (ao)

add row

As an globally seen in O by all processors

Eventuality: read by processor from location X that follows a wr by another processor to X returns Wr if:

① Read & wr sufficiently separated in time ② No other wr occur b/w R & W to X

Full bit vector scheme: ^{Bad scaling} each dir entry has state of cache line & bit vector w/ one shared bit per 4 dir states → dir bits per cache line? ex 128 cores w/ priv \$ → 2 bit ($2^2=4$), 128 bits (bit vec) + 2

Hierarchical bit vector scheme: aggr. cores into grp → top as single bit, invalids sent to all in grp

ex 1024 core 64 byte cache lines, # cores / group to reduce dir state by 10% of phys mem?

$2 + 1024/N \leq 64 \times 8/10 \rightarrow 2 + 1024/N \leq 512/10 \rightarrow N \geq 21 \rightarrow$ at least 21 cores

★ Must store dir bits for every line in mem (un if not \$d) → hierarchal dir struct.

→ or add last level cache (LLC): stores only recently accessed; inclusive of smaller \$

False sharing: 2+ thrds modify data that happens to reside on same \$ line;

inefficient since each thr. may invalidate \$ line for other thrds

ex sharing \$ line w/ single bit; even if val didn't change its force to be I

Mem consistency across all mem addr; even if no caches, important for multi thr.

Sequential Consistency (SC): order-preserving P1: A, B P2: C, D → BADC not allowed

★ Most real time not SC → slower, reads complex HW → reads TSO to work

Store buffer optimization: CPU → St buff → Shared mem; later loads can go ahead of st if diff

Total Store Ordering (TSO): Relaxes W → R (R can be visible before W)

★ allow use of store buff, sd → sd could take ∞ time

Weak-multiprocessor atomic mem: all processors see writes by another processor in same order

(relaxes everything outside but internally everything is in order) "pt of visibility"

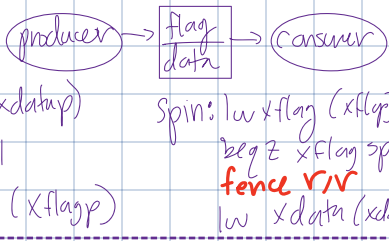
... NON-multiprocessor ...: write can be partially visible to all cores (intermediate buffer)

AKA hierarchal shared buffering, relax all, maintain inner order

Weak Simpler to achieve high perf, easy HW, expose HW to SW Strong easy on ISA, high complexity

Fences

SW xdata (xdatap)
li xflag, 1
fence r/r
sw xflag, (xflagp)



Relaxed mem models:

$X \rightarrow Y$: X must be visible before Y is visible

$R \rightarrow R, R \rightarrow W, W \rightarrow R, W \rightarrow W$

Release consistency: relax all four

★ compilers can reader $\rightarrow$ bad performance

initial state	other cached	ops	Actions by this cache	final state	this cache	other caches	mem
clean Exclusive A CPU wr/rd = self processor issues, write hit/miss	no bus signals	none	none	CE	yes		yes
		CPU read	none	CE	yes		yes
		CPU write	none	OE	yes		
		replace	none	I			yes
		CR	none or CCI'	CS	yes	yes	yes
		CRI	none or CCI'	I		yes	
		CI	none		impossible (if shared)		
		WR	none		impossible (write)		
		CWI	none	I			yes

ID	Event	Message Shared	Cache0 State	Cache1 State	Cache2 State	Main Memory Up-to-Date?
0	CPU0: read A	0:OR	OE	I	I	Yes
1	CPU2: write B	2:CR	I	I	OE	No
2	CPU1: read B	1:CR; 2:CCI	I	OS	OS	No
3	CPU1: write B	1:CI	I	OE	I	No
4	CPU2: write A	2:CR; 1:CCI	I	I	OE	No
5	CPU0: write B	0:CR; 2:CCI	OE	I	I	No

Guest lectures

Apple

- Branch predictions
- Cost of mispredict grows w/ deeper pipes
- Static (simplest)
 - ↳ backward (loops) usually taken
 - ↳ fwd is more unknown
- Dynamic (history)
 - ↳ Bimodal predictor: FSM indexed w/ PC
- Predictor aliasing: 2 diff branches w/ diff PC & diff bias share same predictor index
 - ↳ Agree predictor: BTB + tags
 - ↳ Gslaved: odd # tables; majority vote
 - ↳ Perception (ML)

(#) is TAGS COR

- next line
- Instr stride (use PC to index into PF Table)
- Irregular
 - ↳ correlation
 - ↳ content directed
 - ↳ execution-based
- Hybrid: use mul. PF to cover patterns, but complex BW intensive
- Feedback: throttle

WSC

- WSC: single org, homog, common layers / SVU platforms in-house services, renting VMs, request-level parallelism
- Parts: server, racks, Top of R networking switch, Accelerators (demand for ML) GPUs, TPUs, FPGAs, VCU
- Networking: hard to fix biBW
 - ↳ N parts for N blocks vs. fat tree, Clos @ Jupiter
- Power: transformers, UPS, PDs
- Untempered Power Sys (UPS): generator, remote sys / distant
- WSC cooling: hierarchy of loops
 - Open loop: re-circulate, transfer
 - Close loop: re-circulate, transfer
- Efficiency = $\frac{\text{compute}}{\text{total energy}} = \left(\frac{1}{\text{PUE}}\right) \times \left(\frac{1}{\text{SPUE}}\right) \times \left(\frac{\text{Comp}}{\text{total Energy (comp)}}$
- PUE = facility power / IT equip power
- SPUE = server power conversion efficiency
- (c) server arch eff. (a) facility efficiency

- Fault tolerance:
 - ↳ Power/cooling: hierarchy, scheduler
 - ↳ Net.: redundancy, high
 - ↳ Scheduling: replication, dist

AWS

- Nitro: more cloud server aka virtualize
 - ↳ server CPU not shared, no limit on inst. size, no tradeoff of network/storage/custom instans
- Scale I/O workload w/ high durability & low-latency
- Like migration
- Security
- Scalability (decouple VA functions from HW), easier migration

Preloading Occupancy = $\frac{\text{latency}}{\text{thruput}}$
Thruput = $\frac{\text{window size}}{\text{latency}}$

PF using performance; change loc of \$ insert