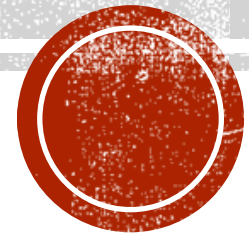


ARCHITECTURE OF COMPUTER SYSTEMS

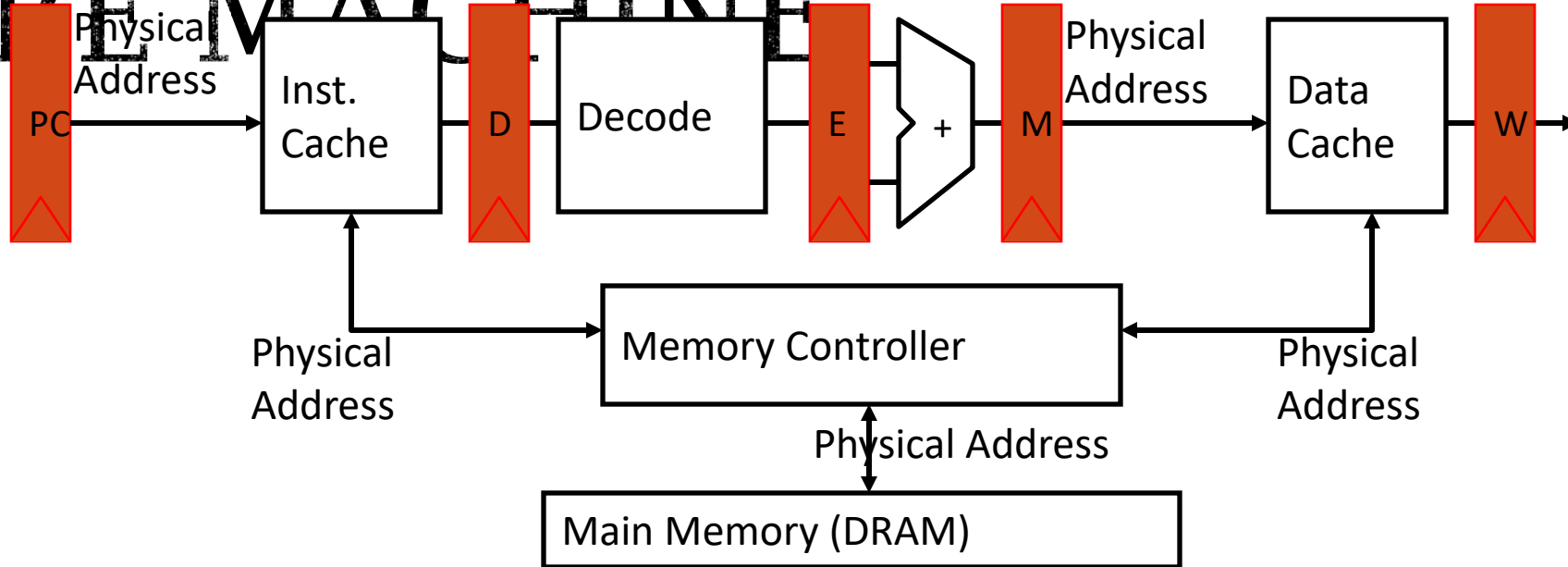
LECTURE 8 - ADDRESS TRANSLATION



LAST TIME IN LECTURE 7

- 3 types of cache misses
 - Compulsory, Capacity, Conflict
- Write policies
 - Write back, write-through, write-allocate, no write allocate
- Multi-level cache hierarchies reduce miss penalty
 - 3 levels common in modern systems (some have 4!)
 - Can change design tradeoffs of L1 cache if known to have L2
- Prefetching: retrieve memory data before CPU request
 - Prefetching can waste bandwidth and cause cache pollution
 - Software vs hardware prefetching
- Software memory hierarchy optimizations
 - Loop interchange, loop fusion, cache tiling

BARE MACHINE



- In a bare machine, the only kind of address is a physical address

ABSOLUTE ADDRESSES

EDSAC, early 50's

- Only one program ran at a time, with unrestricted access to entire machine (RAM + I/O devices)
- Addresses in a program depended upon where the program was to be loaded in memory
- *But* it was more convenient for programmers to write location-independent subroutines

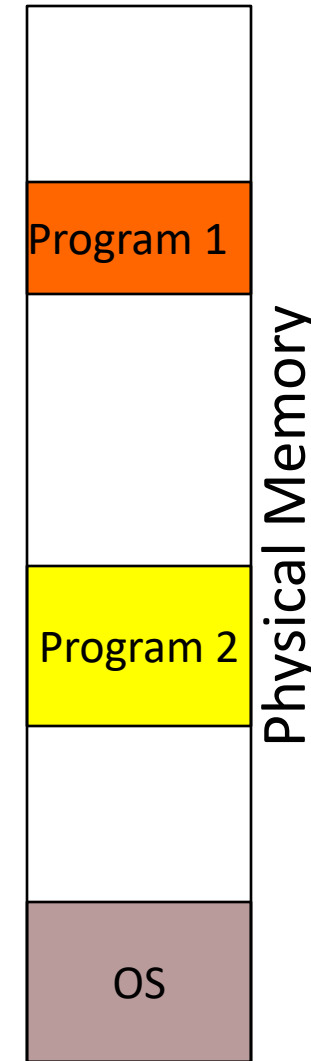
How could location independence be achieved?

Linker and/or loader modify addresses of subroutines and callers when building a program memory image

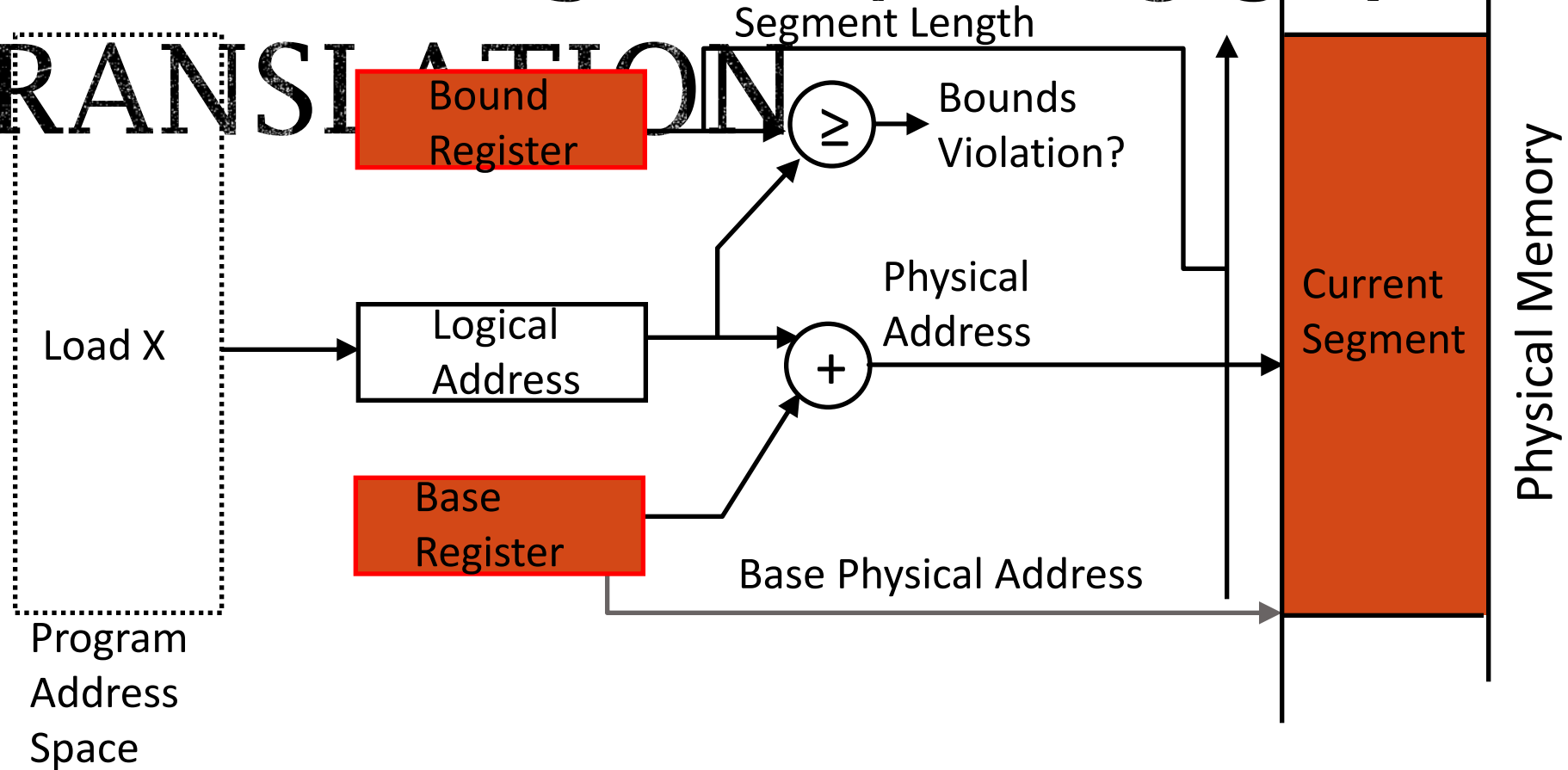
DYNAMIC ADDRESS

TRANSLATION

- Motivation
 - In early machines, I/O was slow and each I/O transfer involved the CPU (programmed I/O)
 - Higher throughput possible if CPU and I/O of 2 or more programs were overlapped, how?
 - multiprogramming with DMA I/O devices, interrupts
- Location-independent programs
 - Programming and storage management ease
 - need for a **base** register
- Protection
 - Independent programs should not affect each other inadvertently
 - need for a **bound** register
- Multiprogramming drives requirement for resident supervisor software to manage context switches between multiple programs



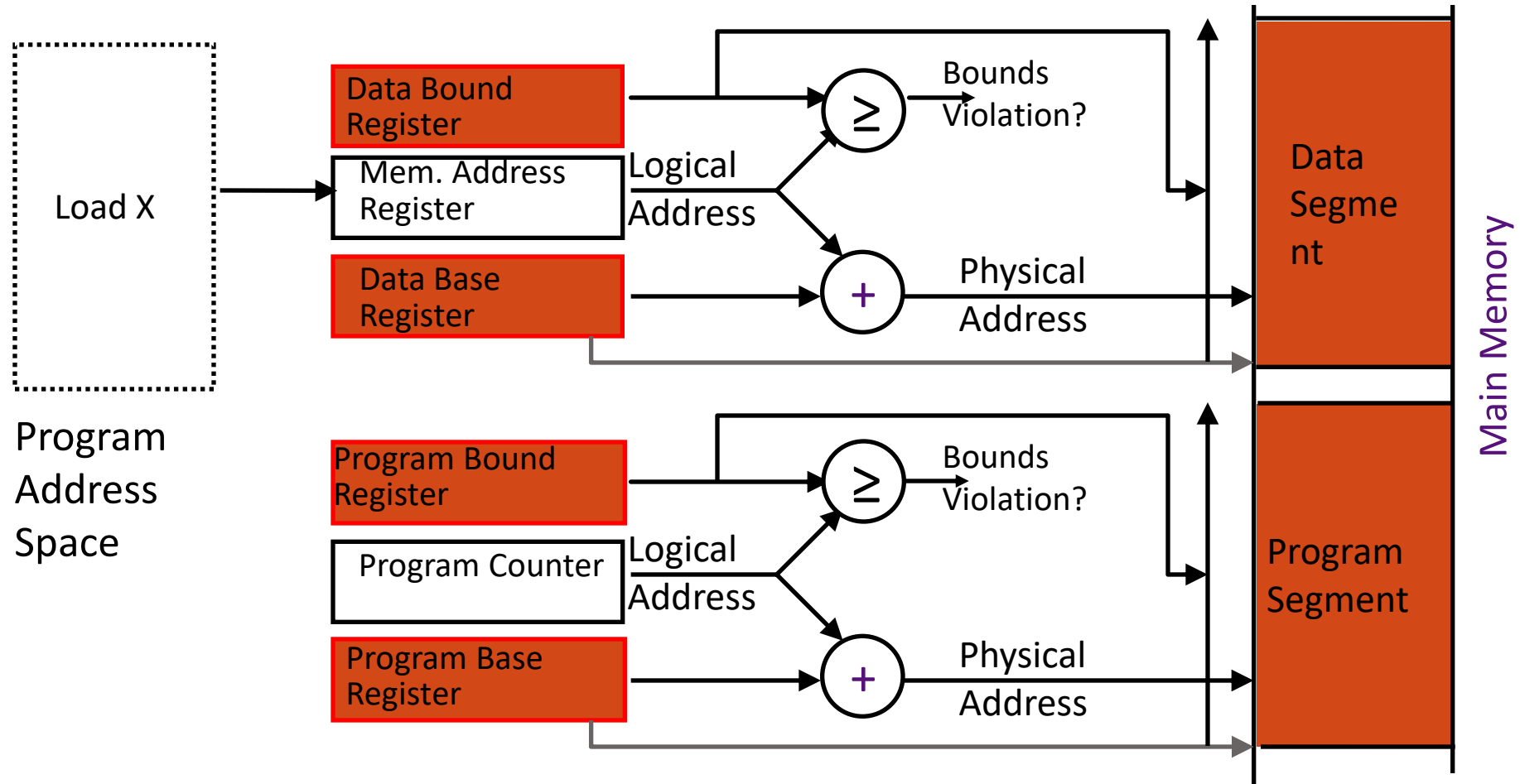
SIMPLE BASE AND BOUND TRANSFORMATION



Base and bounds registers are visible/accessible only when processor is running in the *supervisor mode*

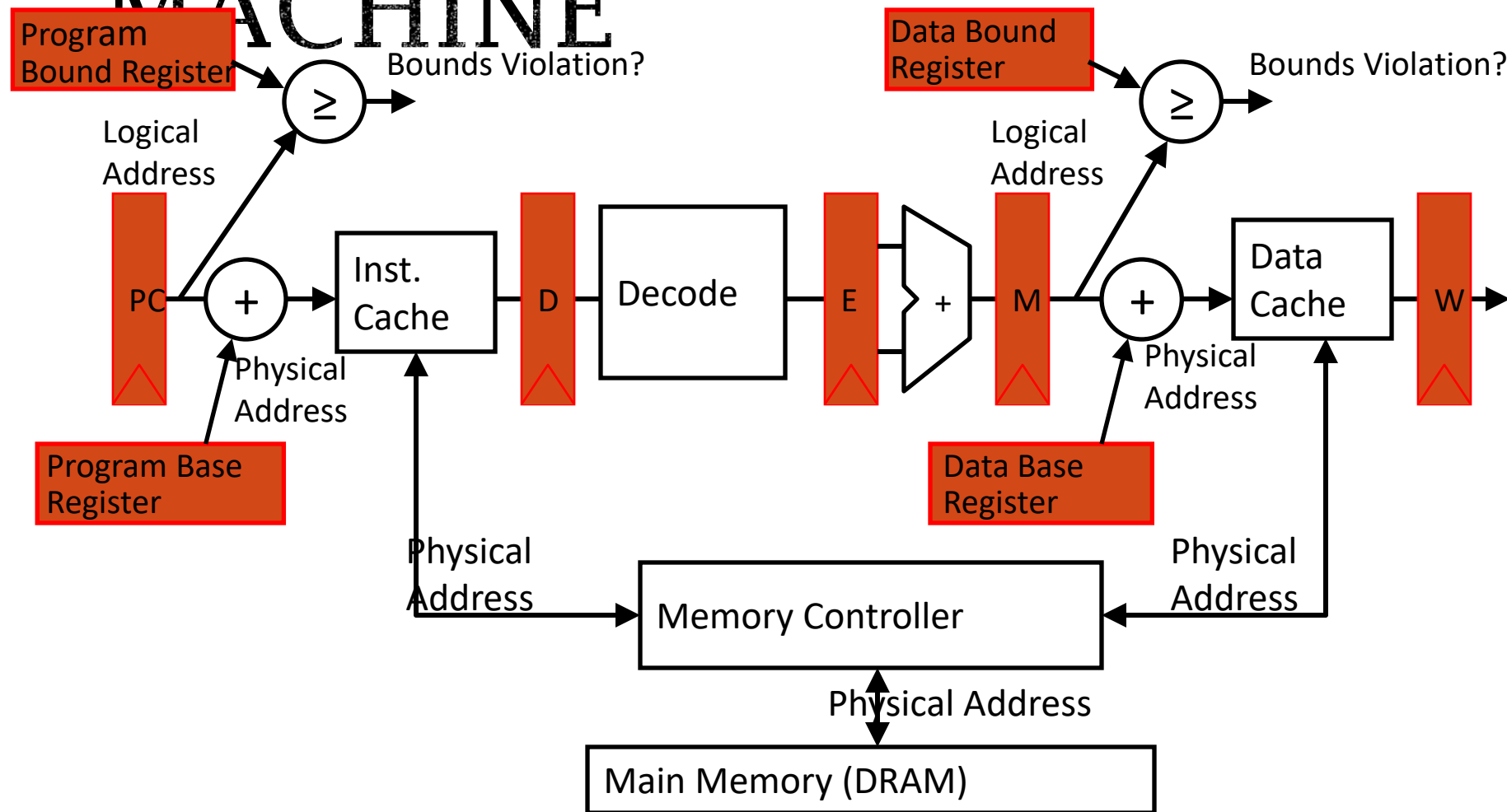
SEPARATE AREAS FOR PROGRAM AND DATA

(Scheme used on all Cray vector supercomputers prior to X1, 2002)



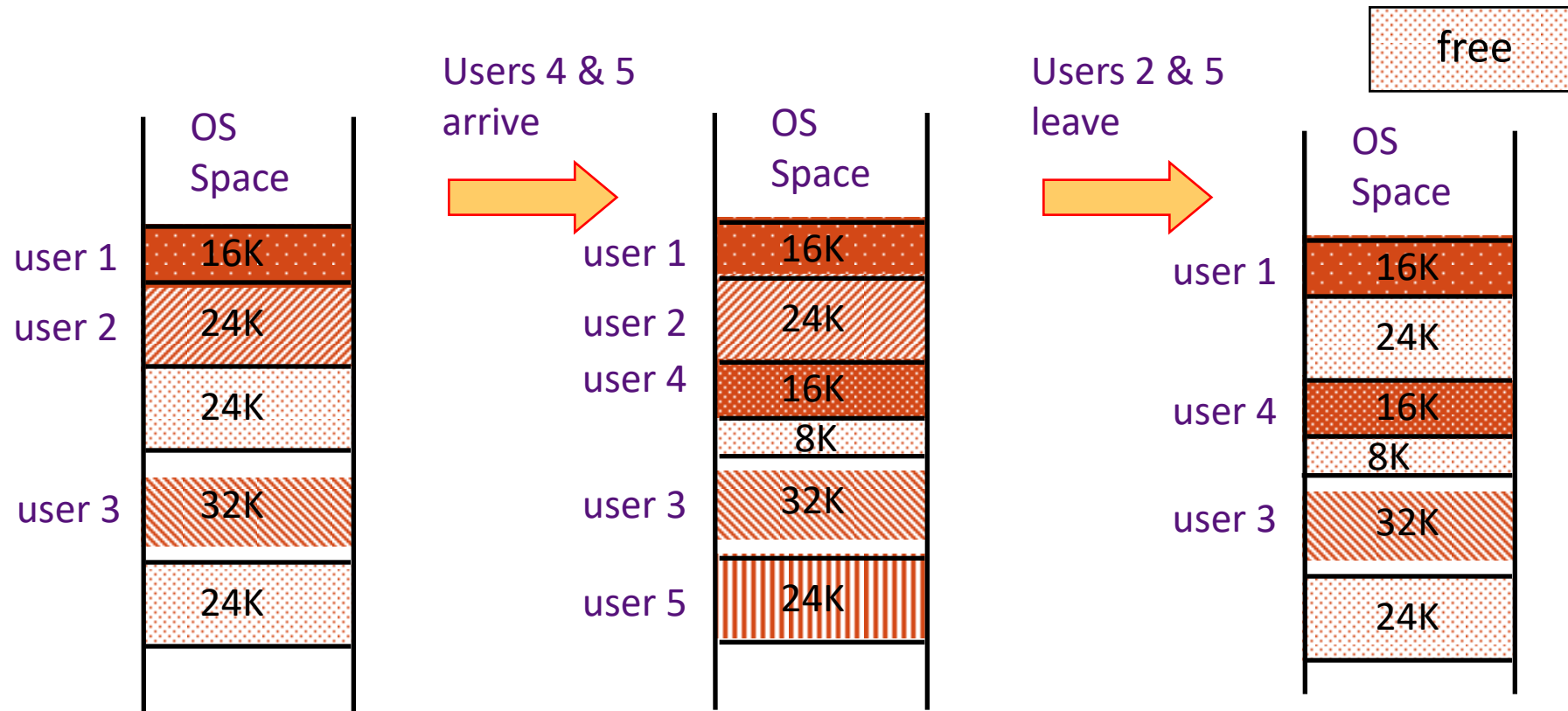
What is an advantage of this separation?

BASE AND BOUND MACHINE



Can fold addition of base register into (register+immediate) address calculation using a carry-save adder (sums three numbers with only a few gate delays more than adding two numbers)

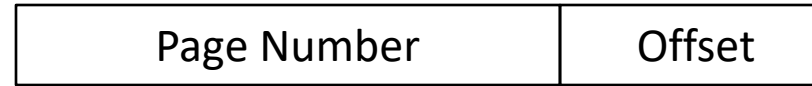
MEMORY FRAGMENTATION



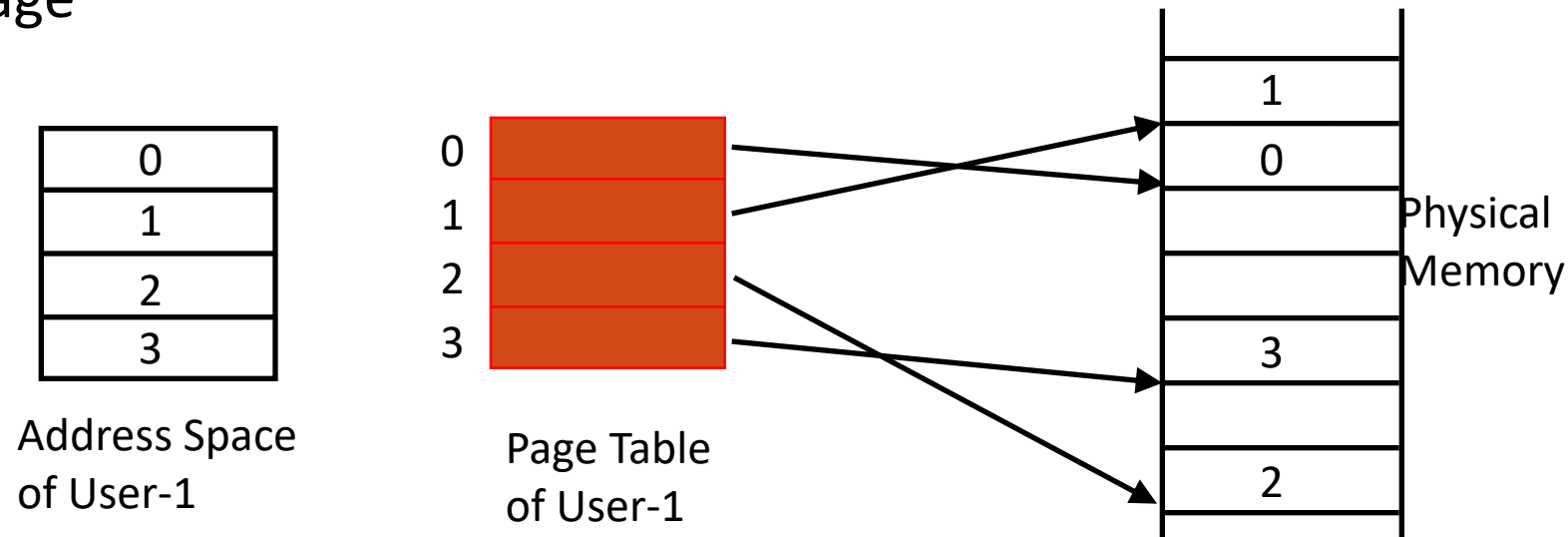
As users come and go, the storage is “fragmented”.
Therefore, at some stage programs have to be moved
around to compact the storage.

PAGED MEMORY SYSTEMS

- Processor-generated address can be split into:

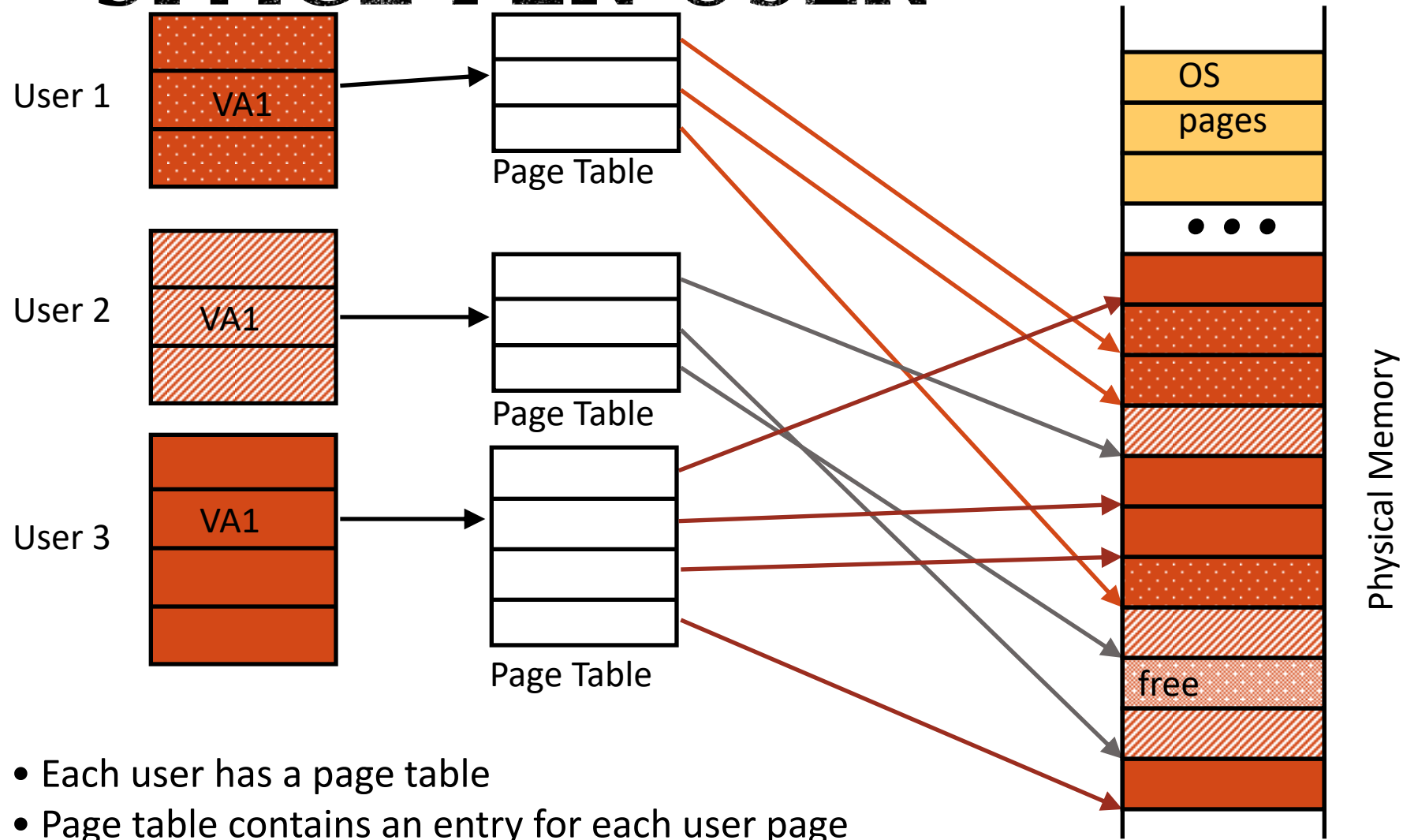


- A Page Table contains the physical address at the start of each page



Page tables make it possible to store the pages of a program non-contiguously.

PRIVATE ADDRESS SPACE PER USER



WHERE SHOULD PAGE

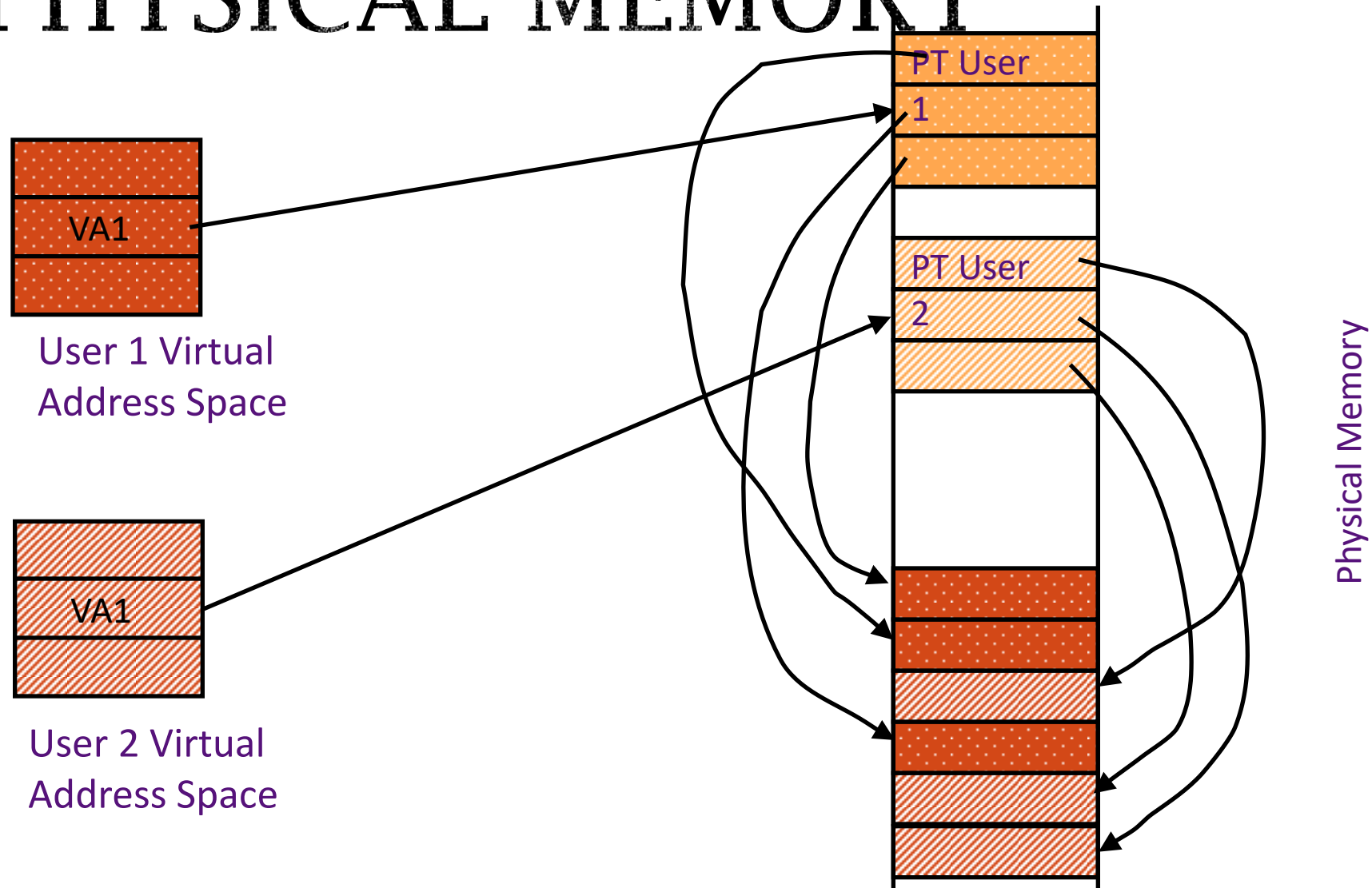
TABLES RESIDE?

- Space required by the page tables (PT) is proportional to the address space, number of users, ...

⇒ Too large to keep in registers

- Idea: Keep PTs in the main memory
 - needs one reference to retrieve the page base address and another to access the data word
- ⇒ doubles the number of memory references!*

PAGE TABLES IN PHYSICAL MEMORY



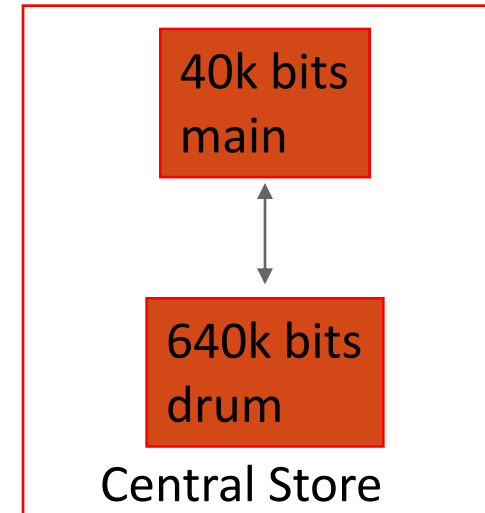
A PROBLEM IN THE EARLY

- ## SIXTIES
- There were many applications whose data could not fit in the main memory, e.g., payroll
 - *Paged memory system reduced fragmentation but still required the whole program to be resident in the main memory*

MANUAL OVERLAYS

- Assume an instruction can address all the storage on the drum

- *Method 1:* programmer keeps track of addresses in the main memory and initiates an I/O transfer when required
 - *Difficult, error-prone!*
- *Method 2:* automatic initiation of I/O transfers by software address translation
 - *Brooker's interpretive coding, 1960*
 - *Inefficient!*



Ferranti Mercury
1956

Not just an ancient black art, e.g., IBM Cell microprocessor using in Playstation-3 has explicitly managed local store!

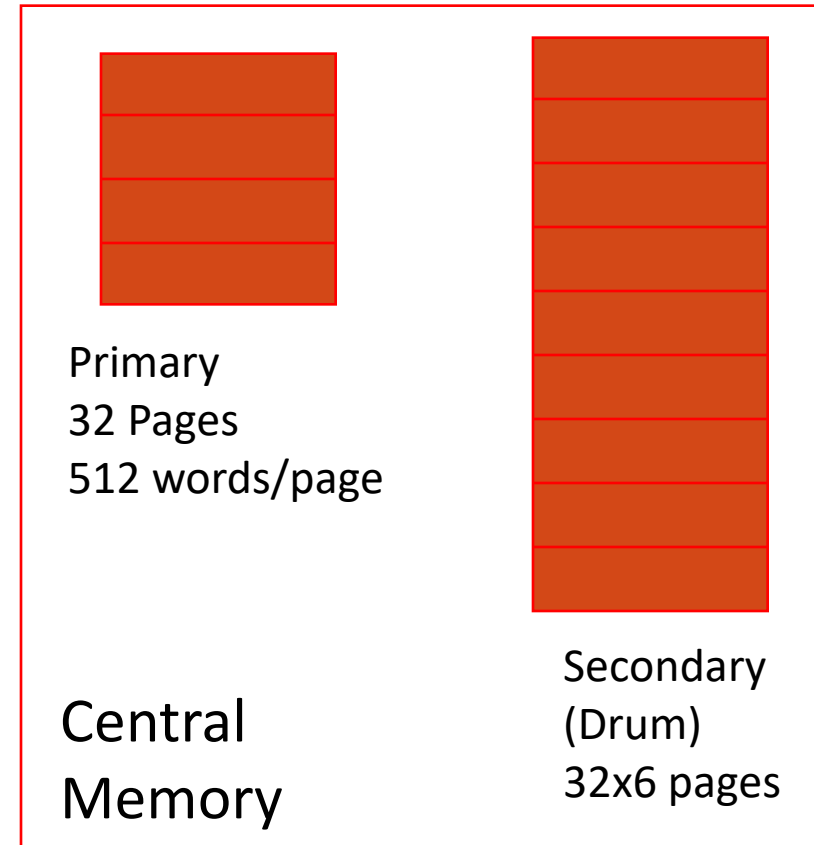
DEMAND PAGING IN ATLAS (1962)

"A page from secondary storage is brought into the primary storage whenever it is (implicitly) demanded by the processor."

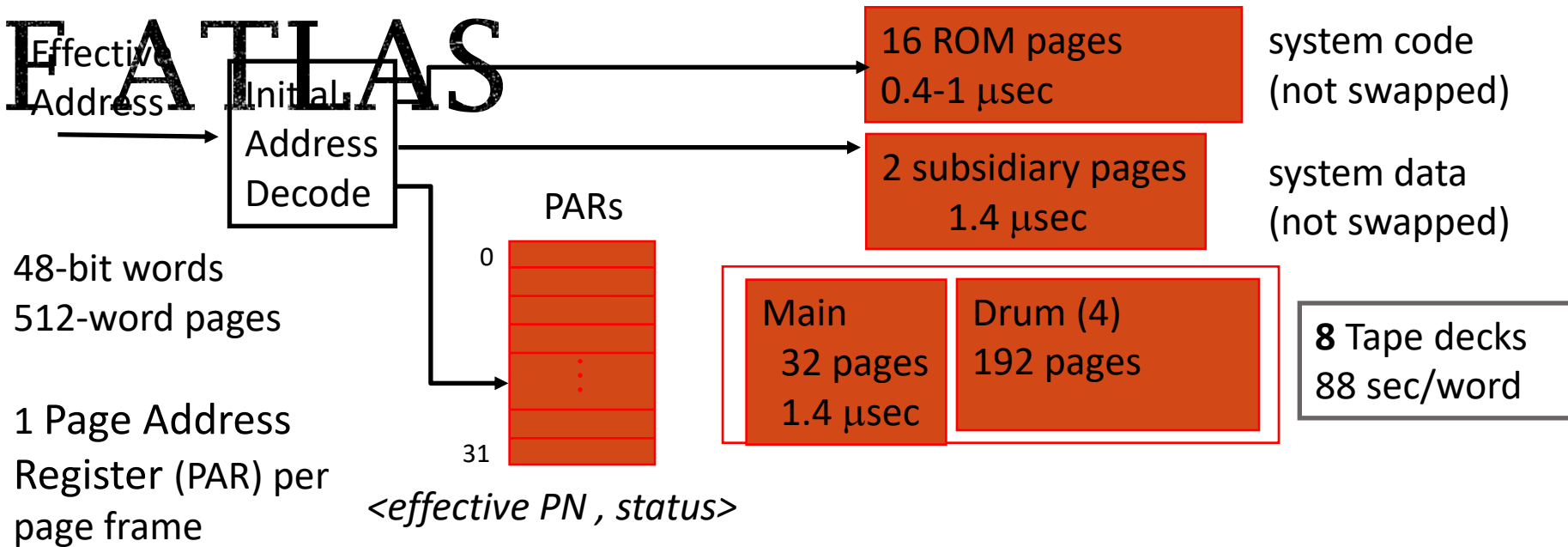
Tom Kilburn

Primary memory as a *cache*
for secondary memory

User sees 32 x 6 x 512 words
of storage



HARDWARE ORGANIZATION OF ATLAS



Compare the effective page address against all 32 PARs

match $\Rightarrow$ normal access

no match $\Rightarrow$ *page fault*

save the state of the partially executed instruction

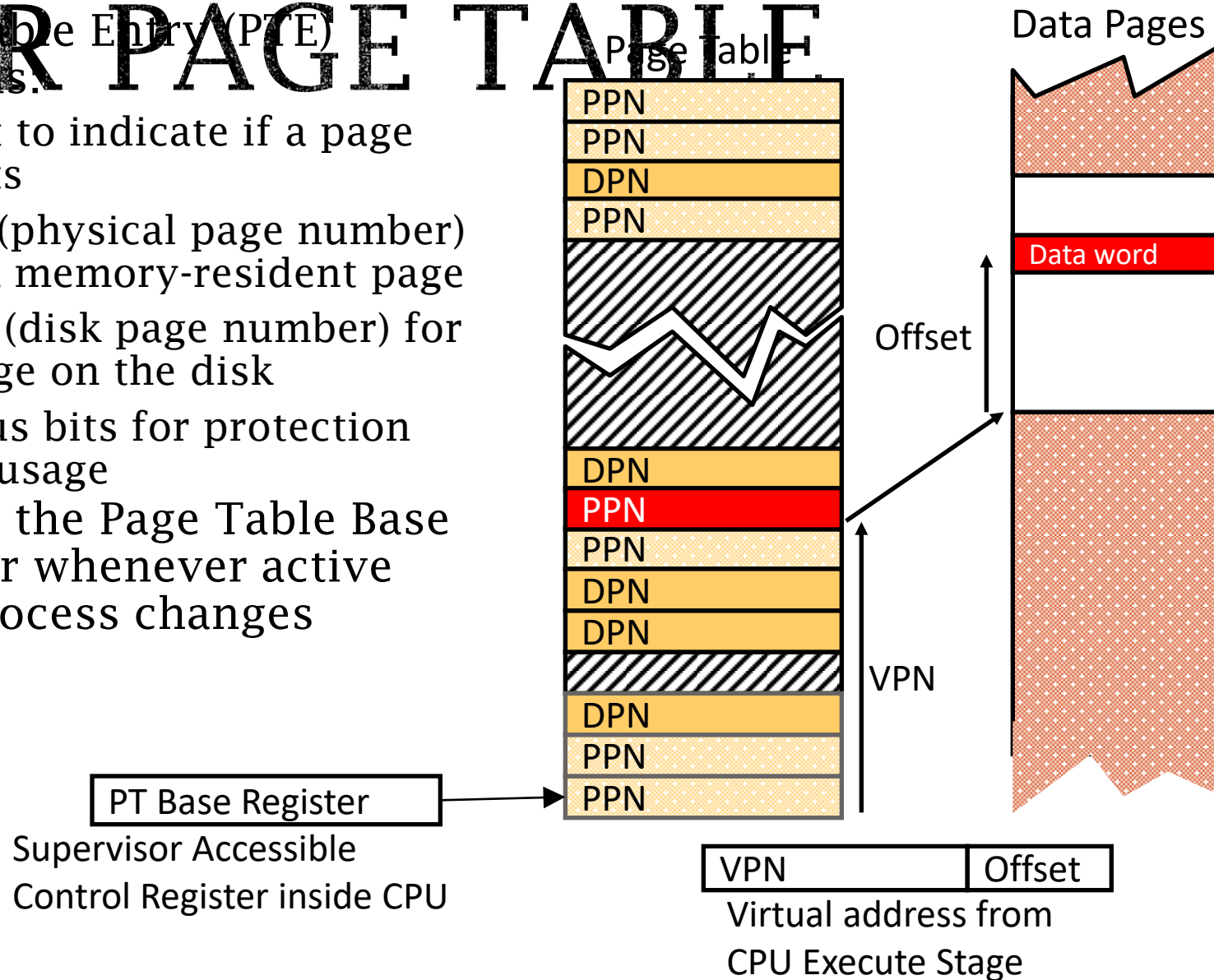
ATLAS DEMAND PAGING SCHEME

On a page fault:

- Input transfer into a free page is initiated
- The Page Address Register (PAR) is updated
- If no free page is left, a page is selected to be replaced (based on usage)
- The replaced page is written on the drum
 - to minimize drum latency effect, the first empty page on the drum was selected
- The page table is updated to point to the new location of the page on the drum

LINEAR PAGE TABLE

- Page Table Entry (PTE) contains:
 - A bit to indicate if a page exists
 - PPN (physical page number) for a memory-resident page
 - DPN (disk page number) for a page on the disk
 - Status bits for protection and usage
- OS sets the Page Table Base Register whenever active user process changes

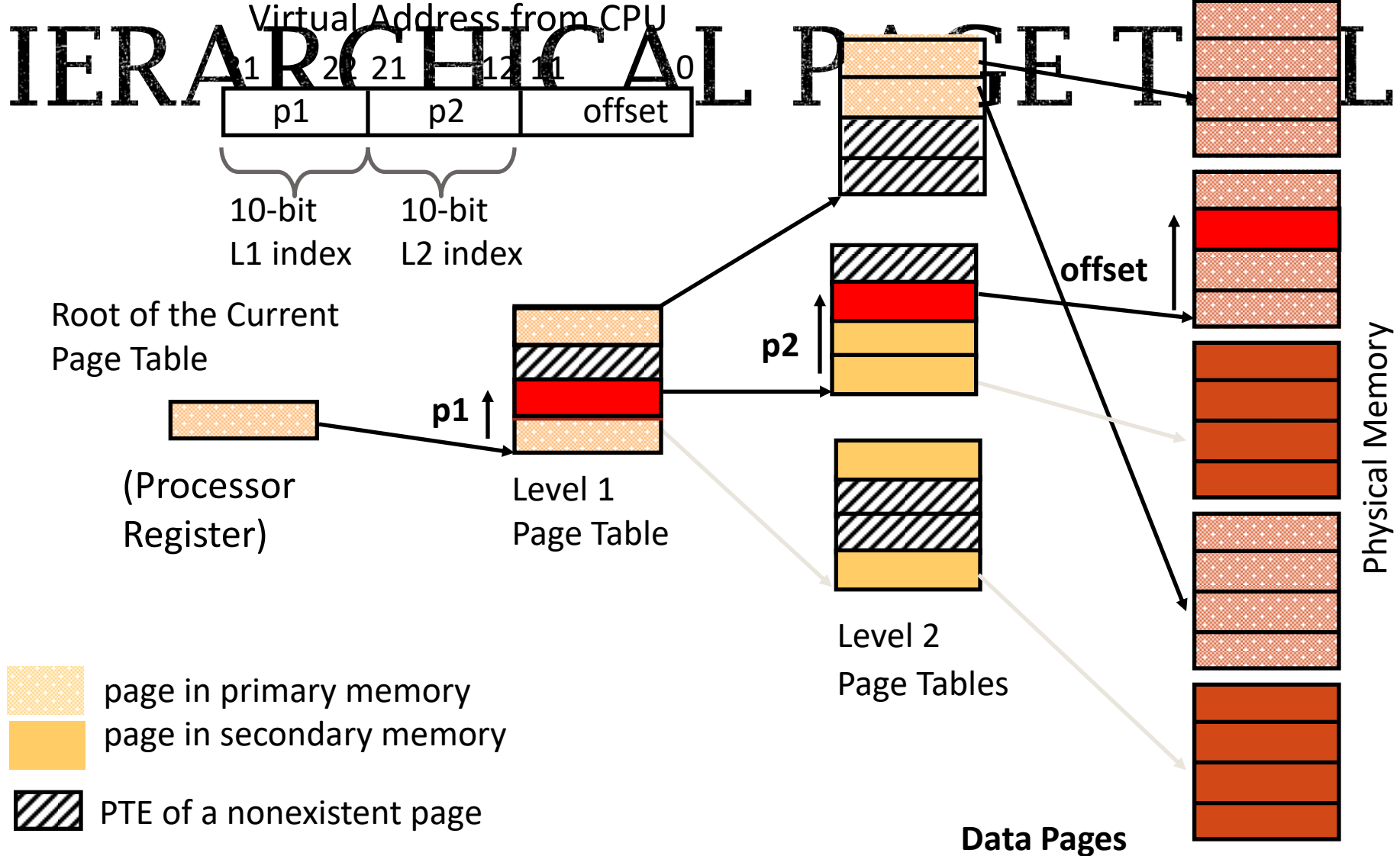


SIZE OF LINEAR PAGE TABLE

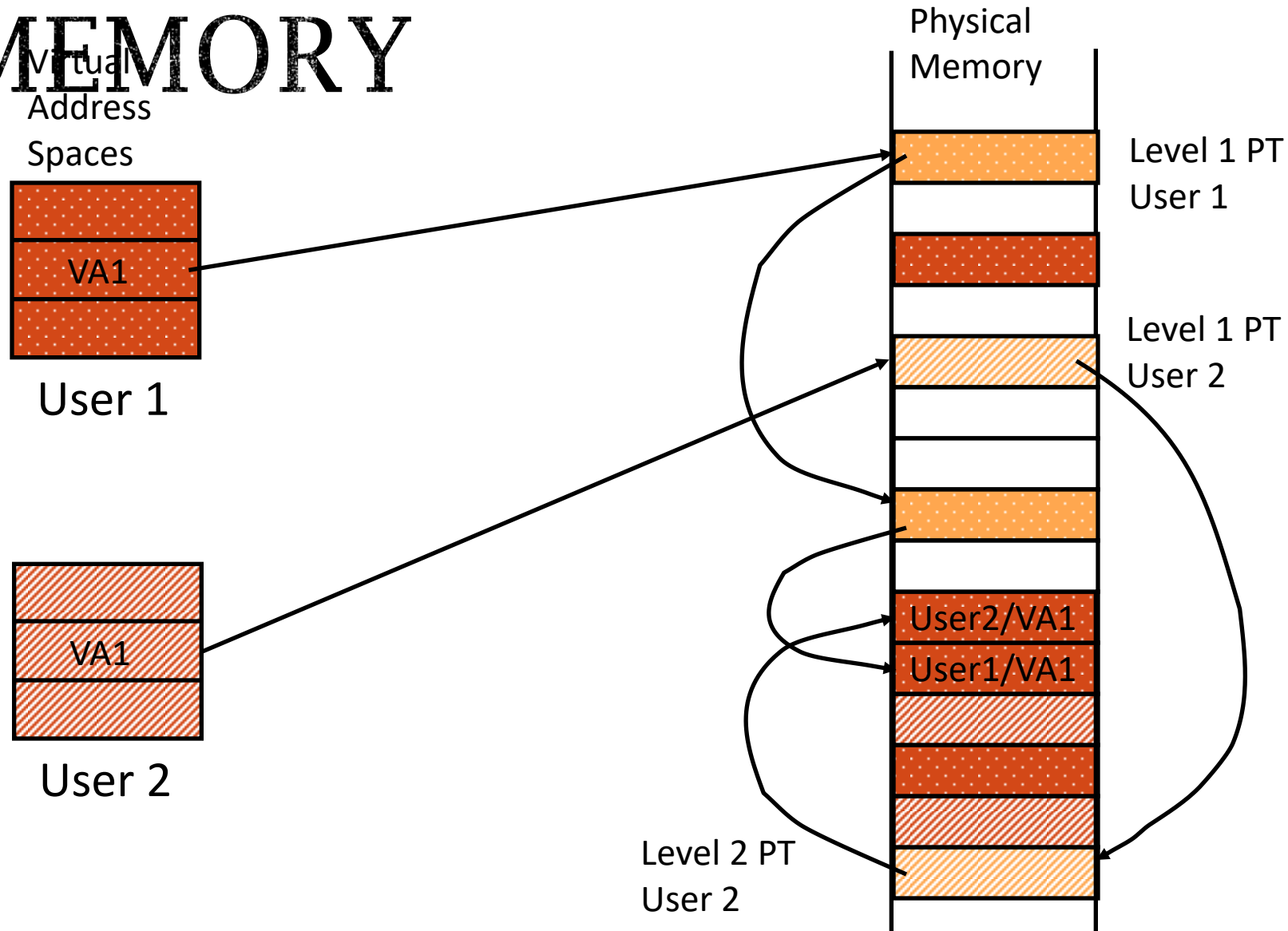
- With 32-bit addresses, 4-KB pages & 4-byte PTEs:
 - 220 PTEs, i.e, 4 MB page table per user
 - 4 GB of swap needed to back up full virtual address space
- Larger pages?
 - Internal fragmentation (Not all memory in page is used)
 - Larger page fault penalty (more time to read from disk)
- What about 64-bit virtual address space???
 - Even 1MB pages would require 2^{44} 8-byte PTEs (35 TB!)

What is the “saving grace” ?

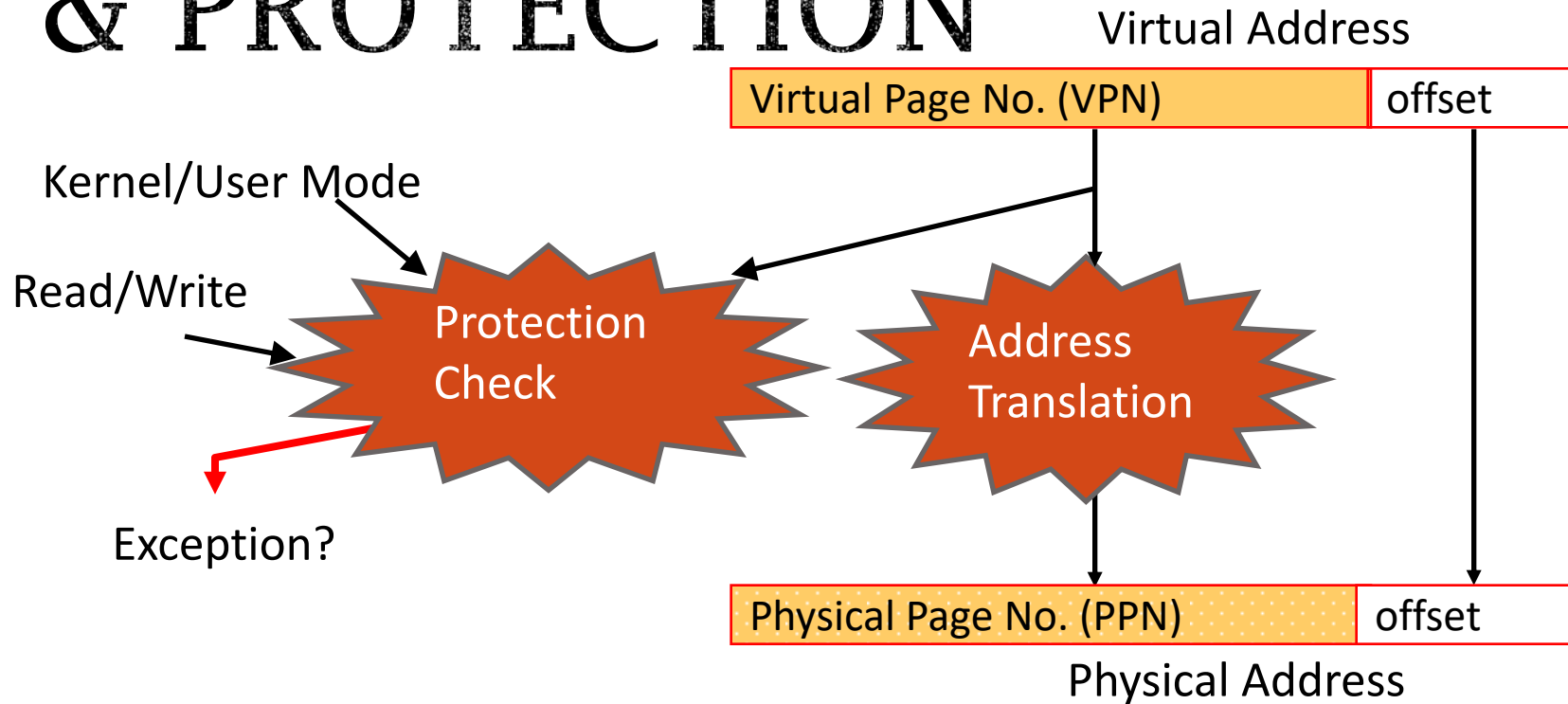
HIERARCHICAL PAGE TABLE



TWO LEVEL PAGE TABLES IN PHYSICAL MEMORY



ADDRESS TRANSLATION & PROTECTION



- Every instruction and data access needs address translation and protection checks

A good VM design needs to be fast (~ one cycle) and space efficient

LOOKASIDE BUFFERS

(TLB)

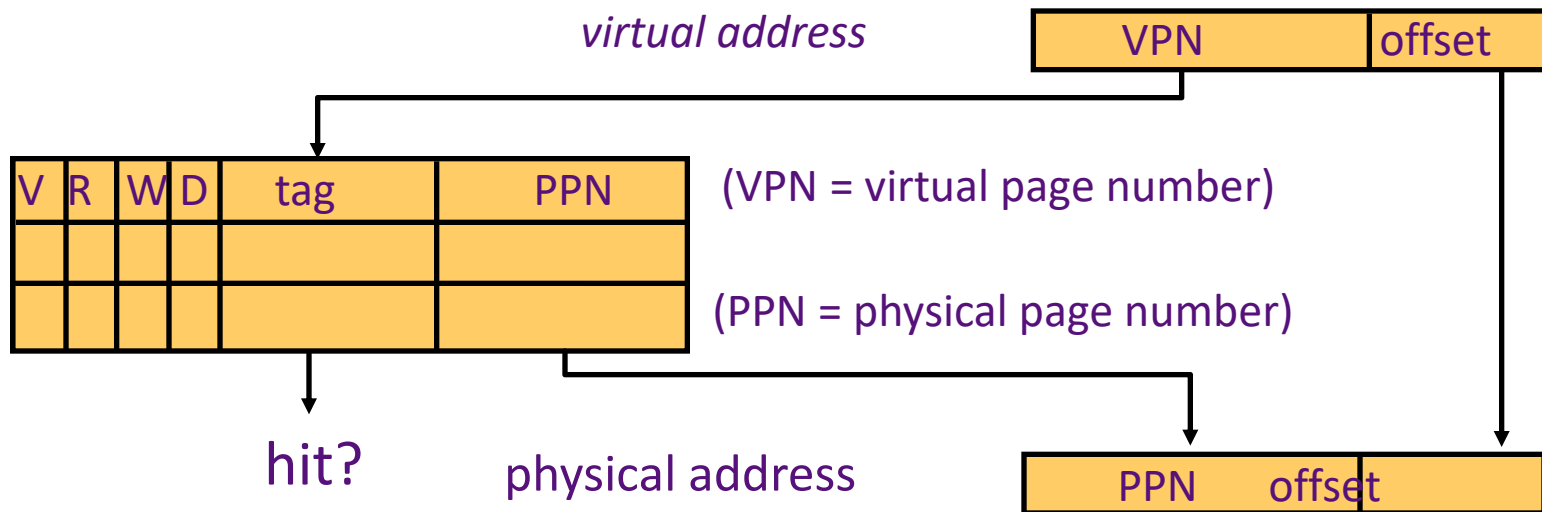
Address translation is very expensive!

In a two-level page table, each reference becomes several memory accesses

Solution: *Cache translations in TLB*

TLB hit $\Rightarrow$ *Single-Cycle Translation*

TLB miss $\Rightarrow$ *Page-Table Walk to refill*



TLB DESIGNS

- Typically 32-128 entries, usually fully associative
 - Each entry maps a large page, hence less spatial locality across pages → more likely that two entries conflict
 - Sometimes larger TLBs (256-512 entries) are 4-8 way set-associative
 - Larger systems sometimes have multi-level (L1 and L2) TLBs
- Random or FIFO replacement policy
- No process information in TLB?
- TLB Reach: Size of largest virtual address space that can be simultaneously mapped by TLB

Example: 64 TLB entries, 4KB pages, one page per entry

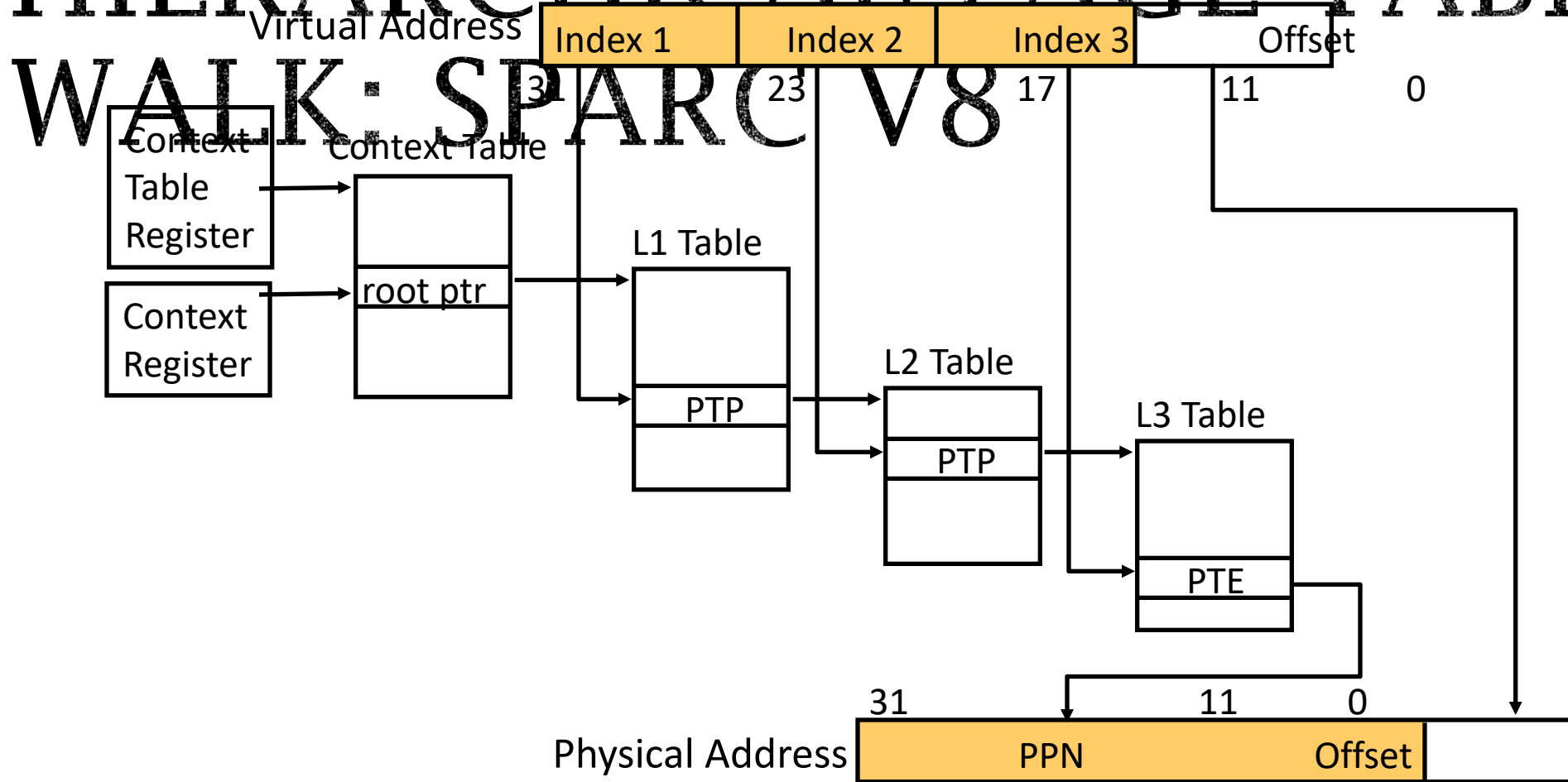
*64 entries * 4 KB = 256 KB (if contiguous)*

TLB Reach = _____?

HANDLING A TLB MISS

- Software (MIPS, Alpha)
 - TLB miss causes an exception and the operating system walks the page tables and reloads TLB. A privileged “untranslated” addressing mode used for walk.
- Hardware (SPARC v8, x86, PowerPC, RISC-V)
 - A memory management unit (MMU) walks the page tables and reloads the TLB.
 - If a missing (data or PT) page is encountered during the TLB reloading, MMU gives up and signals a Page Fault exception for the original instruction.

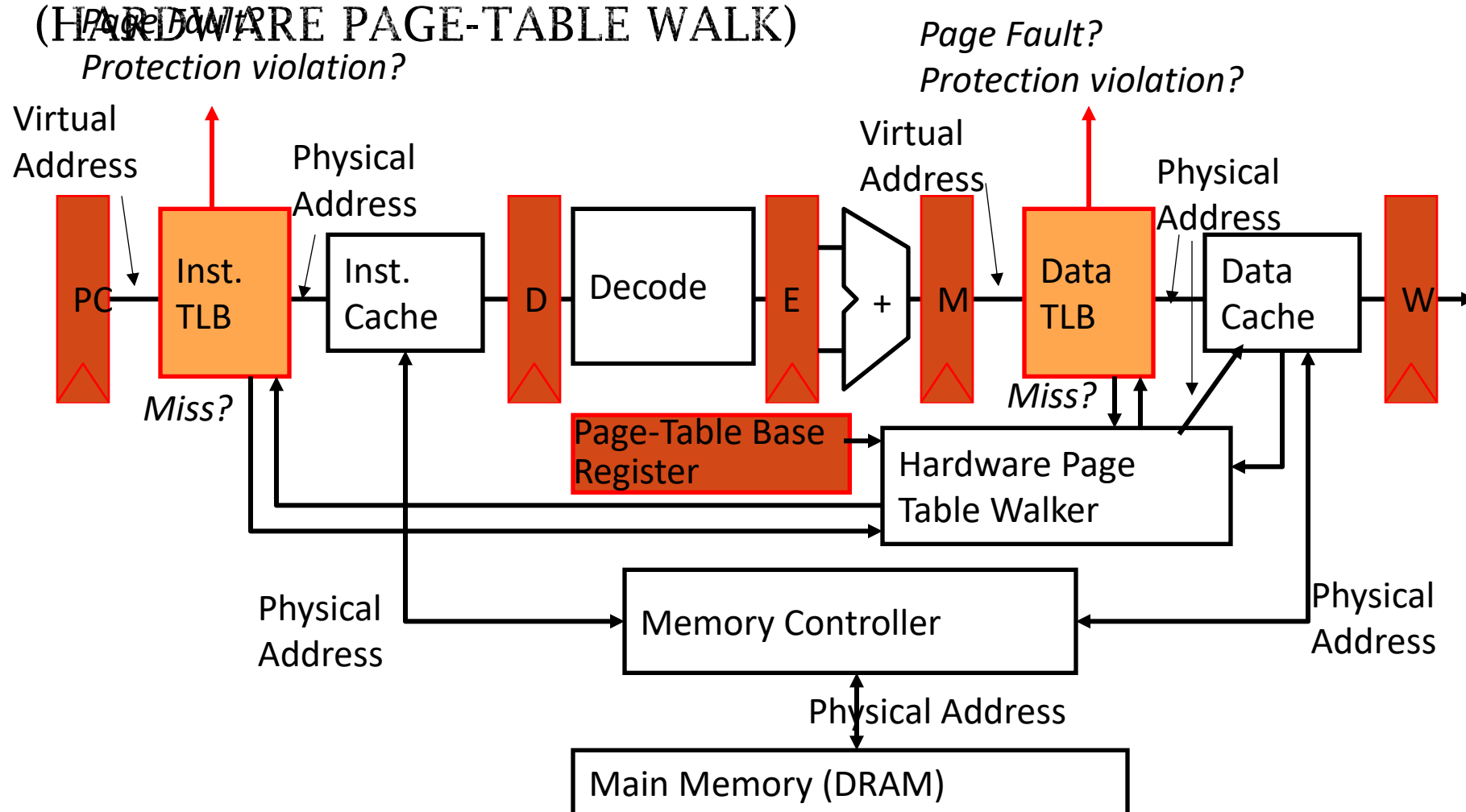
HIERARCHICAL PAGE TABLE



MMU does this table walk in hardware on a TLB miss

PAGE-BASED VIRTUAL- MEMORY MACHINE

(HARDWARE PAGE-TABLE WALK)

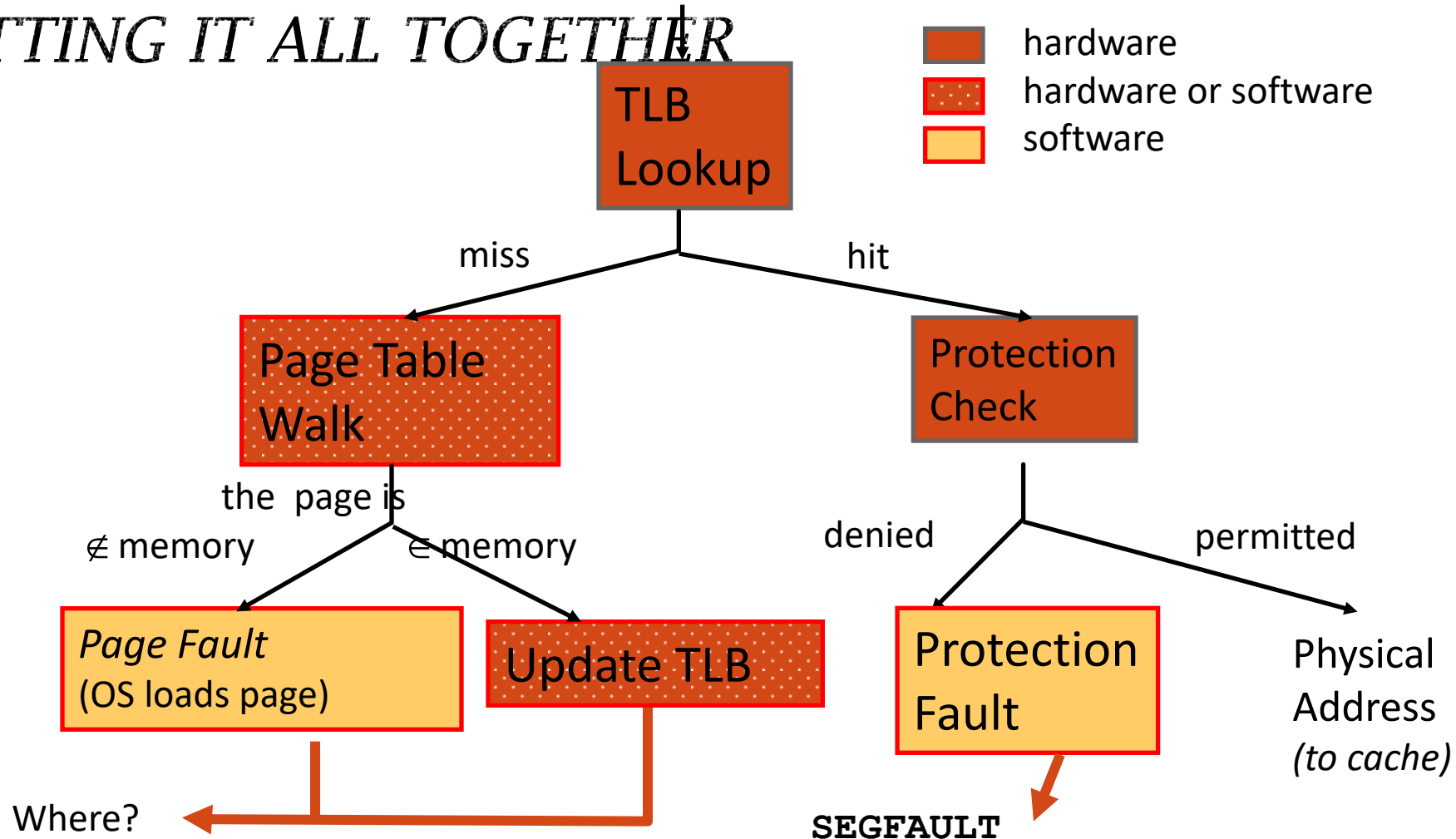


- Assumes page tables held in untranslated physical memory

ADDRESS TRANSLATION:

Virtual Address

PUTTING IT ALL TOGETHER



ACKNOWLEDGEMENTS

- These slides contain material developed and copyright by:
 - Arvind (MIT)
 - Krste Asanovic (MIT/UCB)
 - Joel Emer (Intel/MIT)
 - James Hoe (CMU)
 - John Kubiatowicz (UCB)
 - David Patterson (UCB)
- MIT material derived from course 6.823
- UCB material derived from course CS252