

Memory Management

Learning Outcomes

- Appreciate the need for memory management in operating systems, and the limits of fixed memory allocation schemes.
- Understand fragmentation in dynamic memory allocation, and understand basic dynamic allocation approaches.
- Understand how program memory addresses relate to physical memory addresses, memory management in base-limit machines, and swapping
- An overview of virtual memory management.

Recap: A Process

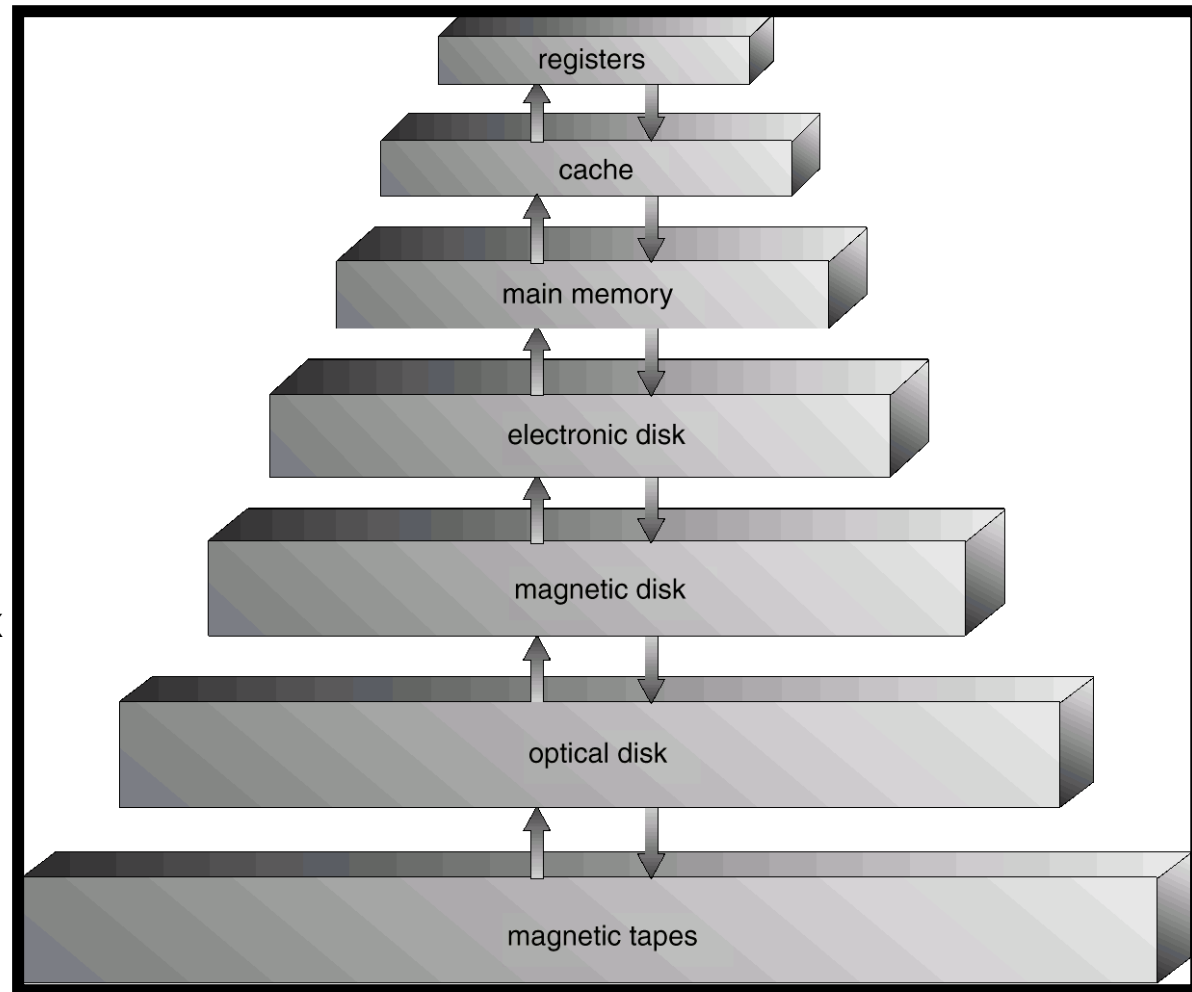
- One or more threads of execution
- Resources required for execution
 - Memory (RAM)
 - Program code (“text”)
 - Data (initialised, uninitialised, stack)
 - Buffers held in the kernel on behalf of the process
 - Others
 - CPU time
 - Files, disk space, printers, etc.

OS Memory Management

- Keeps track of what memory is in use and what memory is free.
- Allocates free memory to processes when needed.
 - Also tracks deallocated memory.
- Manages the transfer of memory content between RAM and disks.

Recap: Memory Hierarchy

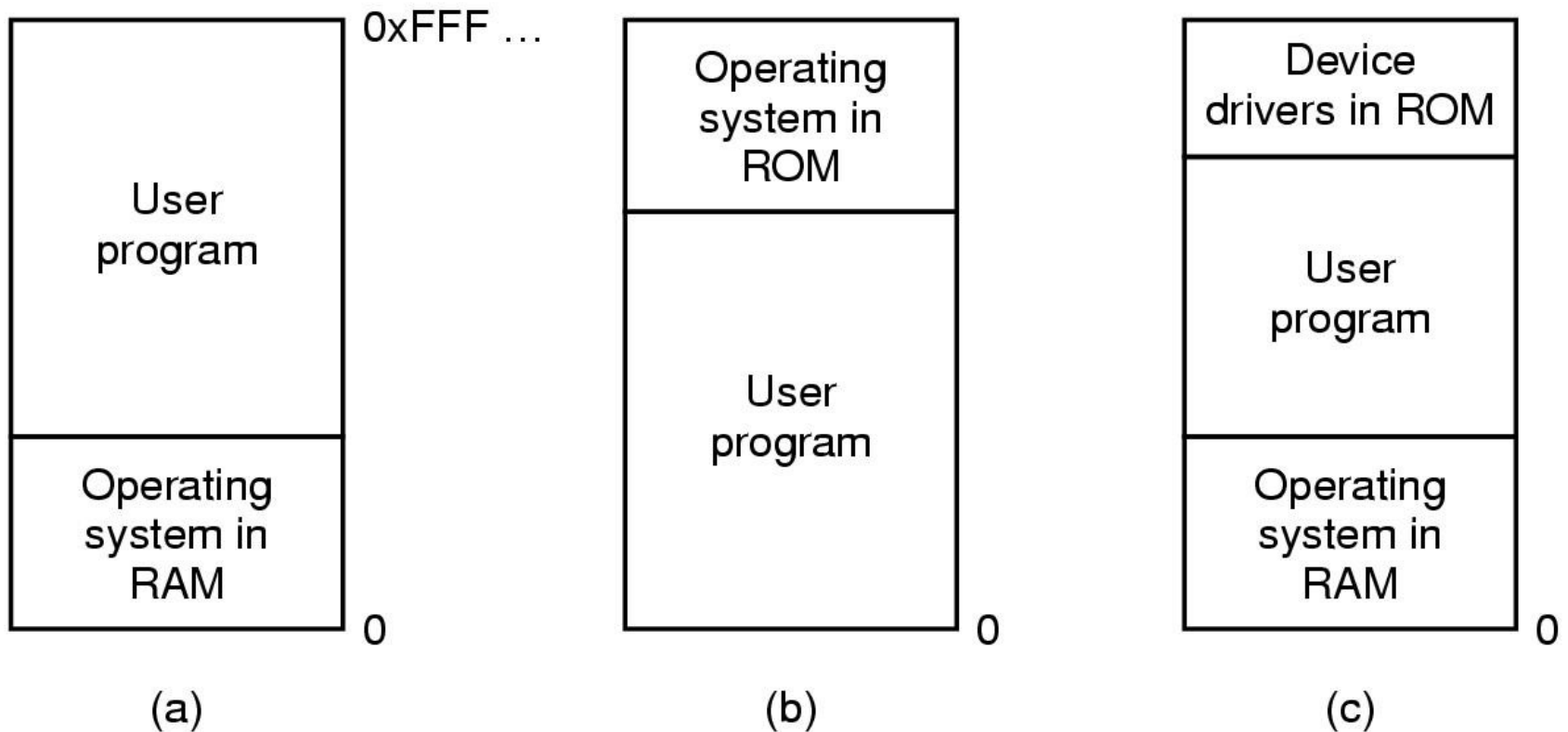
- Ideally, programmers want memory that is
 - Fast
 - Large
 - Nonvolatile
- Not possible
- Memory management coordinates how memory hierarchy is used.
 - Focus usually on RAM $\Leftrightarrow$ Disk



OS Memory Management

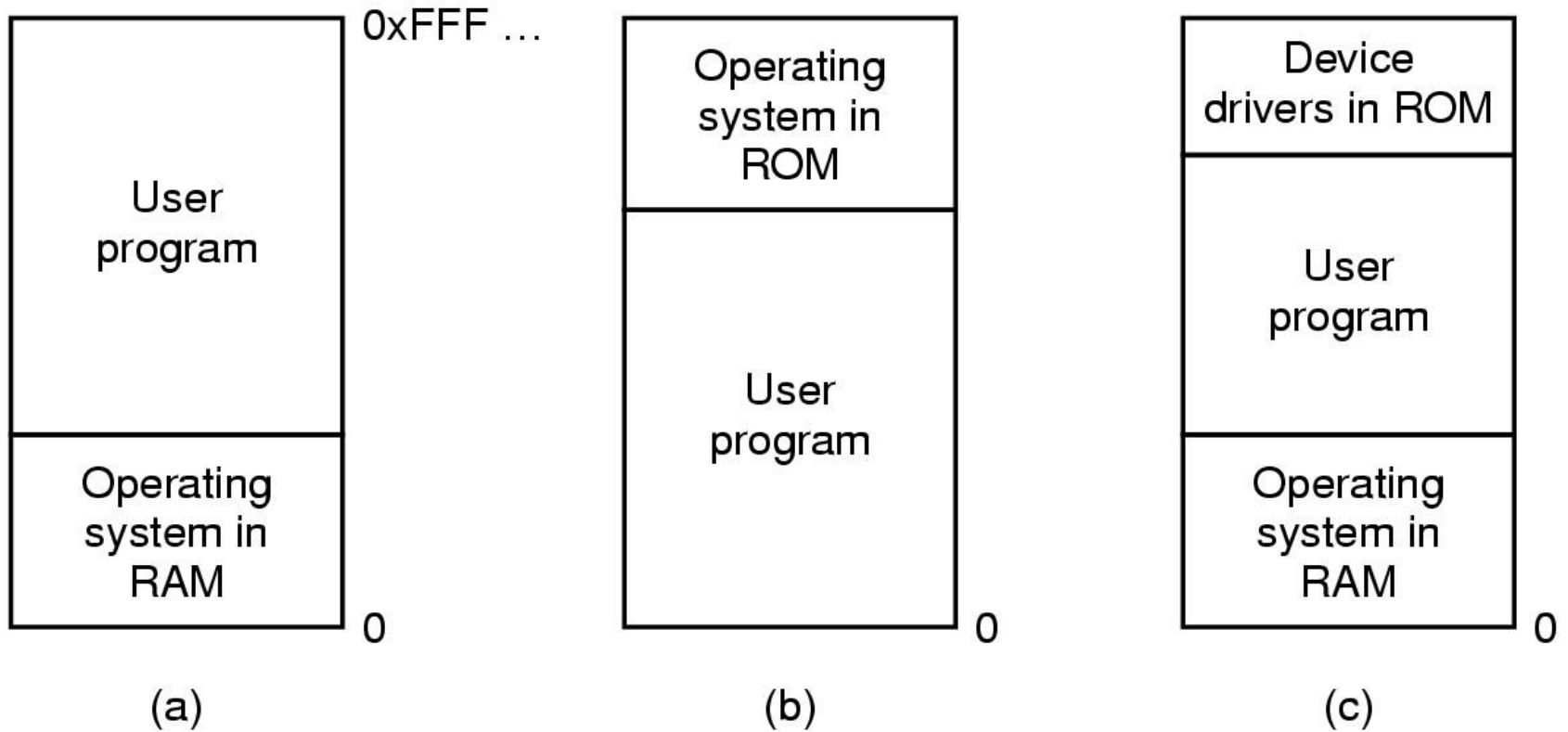
- There are two broad classes of OS memory management systems
 - OSs that transfer processes to and from external storage during execution.
 - Called swapping or paging
 - OSs that don't
 - Simple static OSs
 - Might find this scheme in an embedded device, dumb phone, or smartcard.
- Like other topics in this course, let's start by considering and ruling out the simple approaches

Static OS Memory Layouts



Three simple ways of organizing memory

- an operating system with one user process



These concepts persist:

- Firmware
- BIOS

Monoprogramming

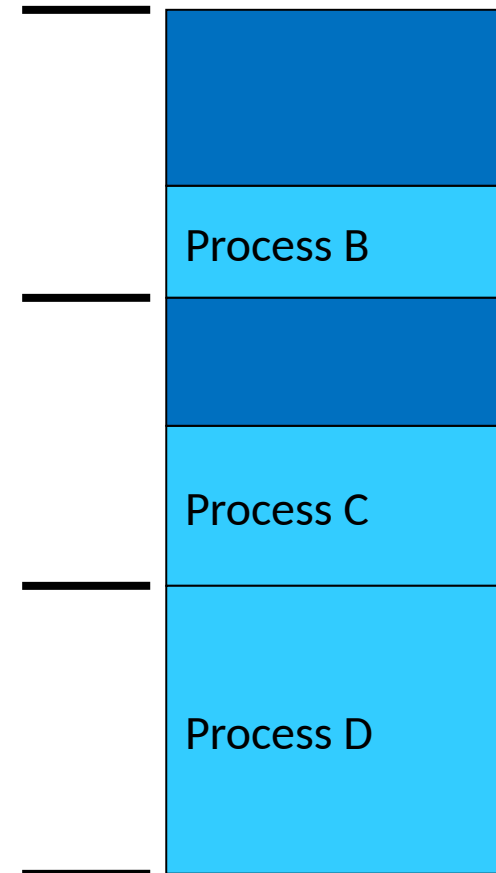
- Monoprogramming = running just one task at a time
- Okay if
 - Only have one thing to do
 - Memory available approximately equates to memory required
- Otherwise,
 - Poor CPU utilisation in the presence of I/O waiting
 - Poor utilisation of memory if we switch between jobs with different memory needs

Idea

- Recall, an OS aims to
 - Maximise memory utilisation
 - Maximise CPU utilization
 - (ignore battery/power-management issues)
- Subdivide memory and run more than one process at once!!!!
 - Multiprogramming, Multitasking

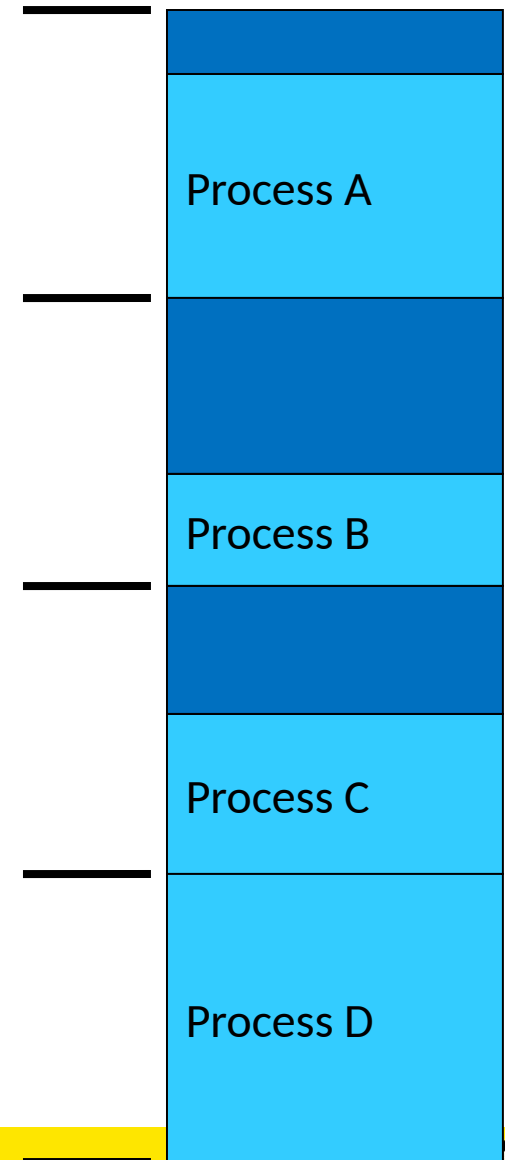
General problem: How to divide memory between processes?

- Given a workload, how to we
 - Keep track of free memory?
 - Locate free memory for a new process?
- Overview of evolution of simple memory management
 - Static (fixed partitioning) approaches
 - Simple, predicable workloads of early computing
 - Dynamic (partitioning) approaches
 - More flexible computing as compute power and complexity increased.
- Introduce virtual memory
 - ~~Segmentation~~ and paging



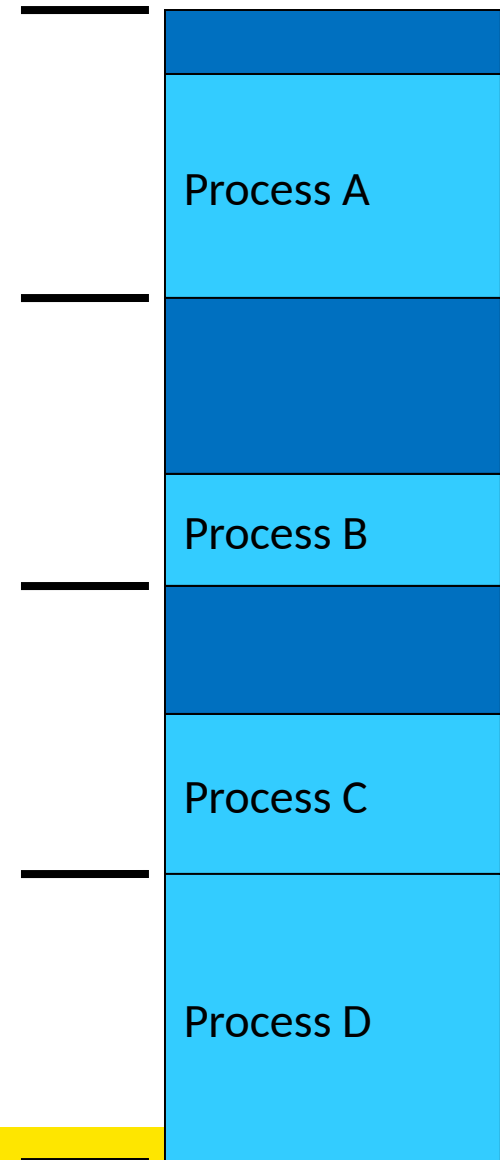
Problem: How to divide memory

- One approach
 - divide memory into fixed equal-sized partitions
 - Any process $\leq$ partition size can be loaded into any partition
 - Partitions are free or busy



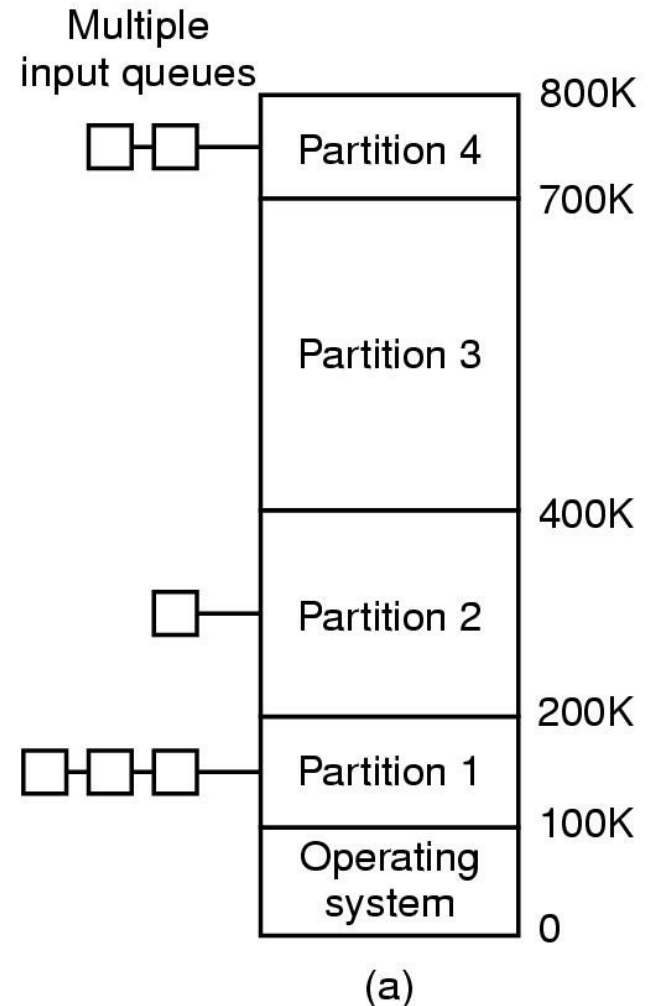
Simple MM (1)

- Simple approach 1: Split memory into fixed, equal-sized partitions.
- Any unused space in the partition is wasted.
 - Called internal fragmentation
- Processes smaller than main memory, but larger than a partition cannot run.



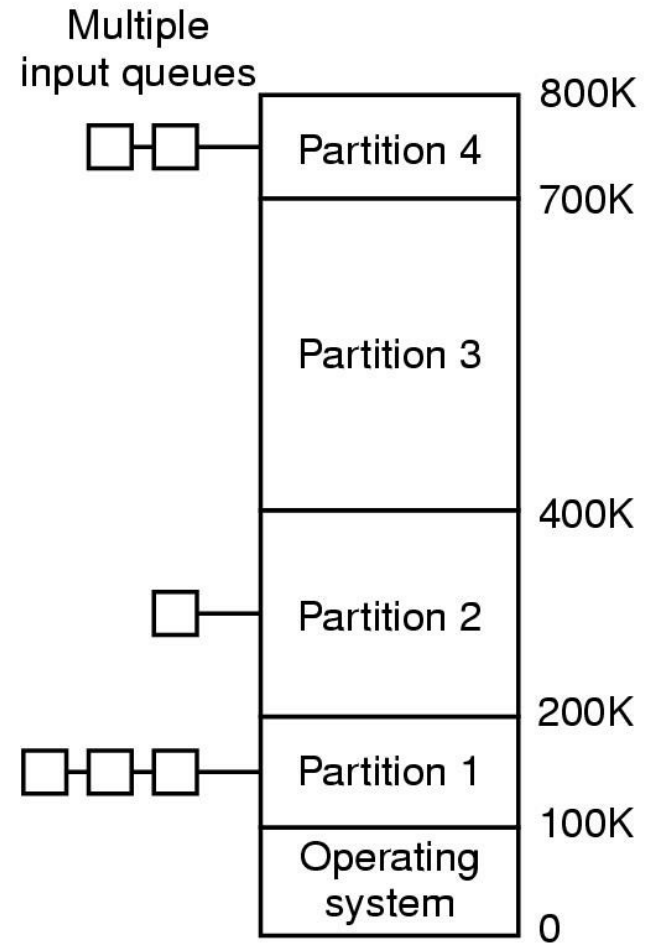
Simple MM (2)

- Simple approach 2: fixed, variable-size partitions
- Divide memory at boot time into a selection of different sized partitions
 - Can base sizes on expected workload
- Each partition has queue:
 - Place process in queue for smallest partition that it fits in.
 - Processes wait for when assigned partition is empty to start



- Issue

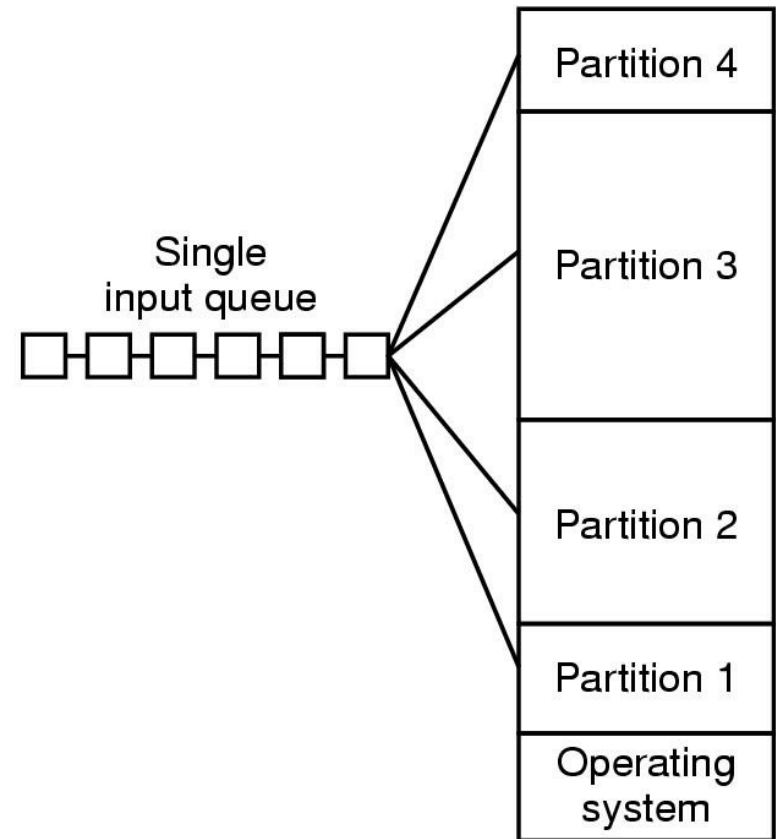
- Some partitions may be idle
 - Small jobs available, but only large partition free
 - Workload could be unpredictable



(a)

Alternative queue strategy

- Single queue, search for any jobs that fit
 - Small jobs in large partition if necessary
- Increases internal memory fragmentation



(b)

Fixed Partition Summary

- Simple
- Easy to implement
- Can result in poor memory utilisation
 - Due to internal fragmentation
- Used on IBM System 360 operating system (OS/MFT)
 - Announced 6 April, 1964
- Still applicable for simple embedded systems
 - Static workload known in advance

Dynamic Partitioning

- Partitions are of variable length
 - Allocated on-demand from ranges of free memory
- Process is allocated exactly what it needs
 - Assumes a process knows what it needs

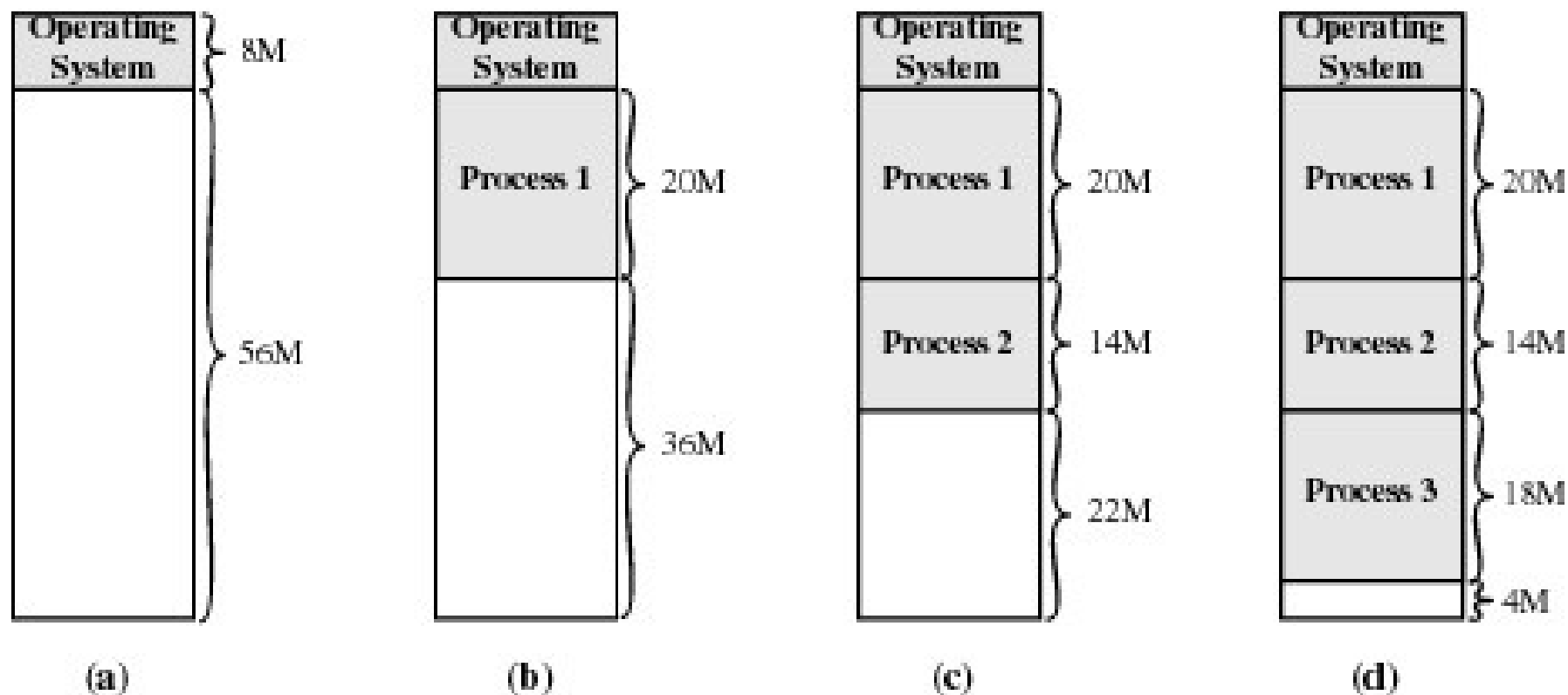


Figure 7.4 The Effect of Dynamic Partitioning

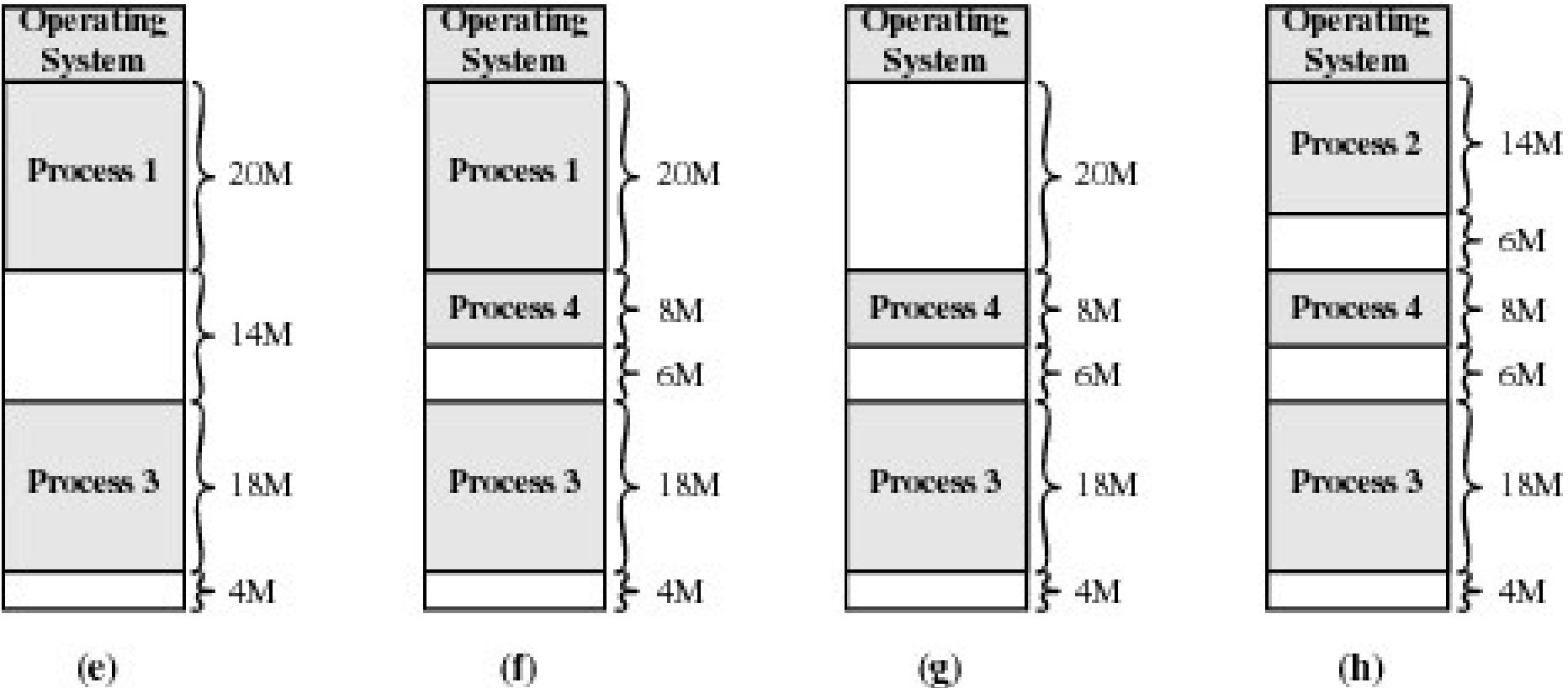


Figure 7.4 The Effect of Dynamic Partitioning

Dynamic Partitioning

- In previous diagram
 - We have 16 meg free in total, but it can't be used to run any more processes requiring > 6 meg as it is fragmented
 - Called *external fragmentation*
- We end up with unusable spaces
 - “fragments”

Recap: Fragmentation

- **External Fragmentation:**

- The space wasted external to the allocated memory regions.
- Memory space exists to satisfy a request, but it is unusable as it is not contiguous.

- **Internal Fragmentation:**

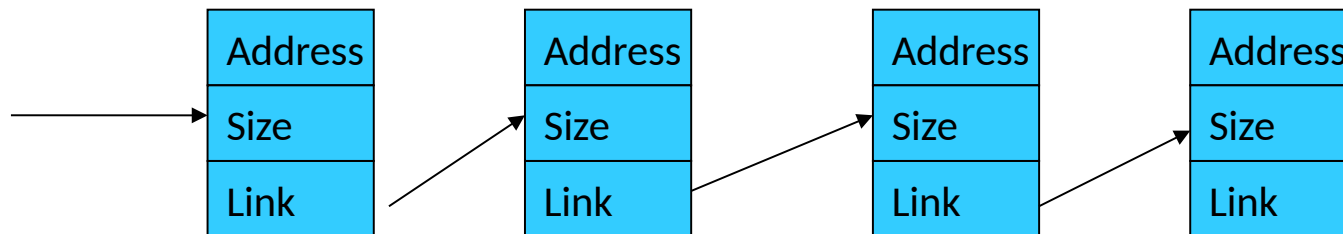
- The space wasted internal to the allocated memory regions.
 - allocated memory may be slightly larger than requested memory; this size difference is wasted memory internal to a partition.
- The same issue came up with file-systems.

Dynamic Partitioning: Allocation

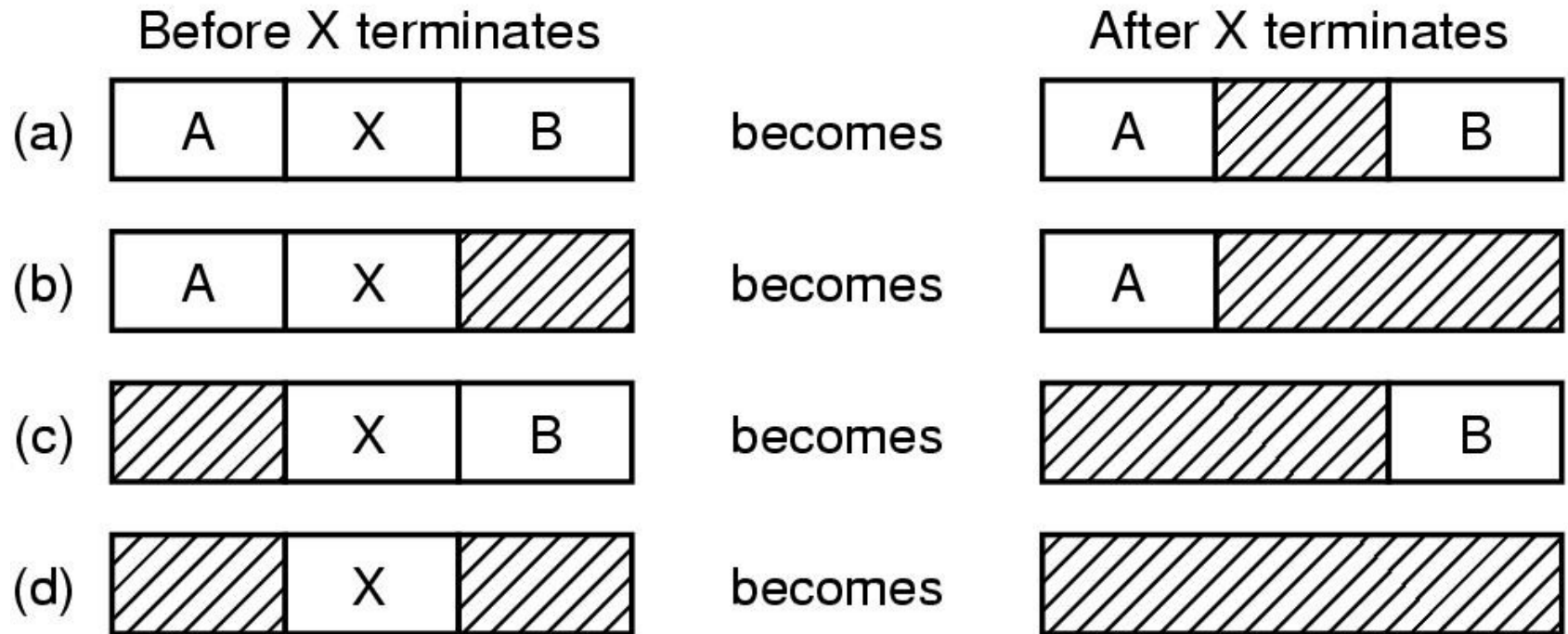
- Dynamic contiguous partition allocation: various free-memory management strategies
- Also applicable to `malloc()` - like in-application allocators
- Given a region of memory, basic requirements are:
 - Quickly locate a free partition satisfying the request
 - Minimise CPU time searching
 - Minimise external fragmentation
 - Minimise memory overhead of bookkeeping
 - Efficiently support merging two adjacent free partitions into a larger partition

Classic Approach

- Represent available memory as a linked list of available “spaces” (free memory ranges).
 - Base, size
 - Kept in order of increasing address
 - Simplifies merging of adjacent spaces into larger spaces.
 - List nodes can be stored in the “spaces” themselves



Coalescing

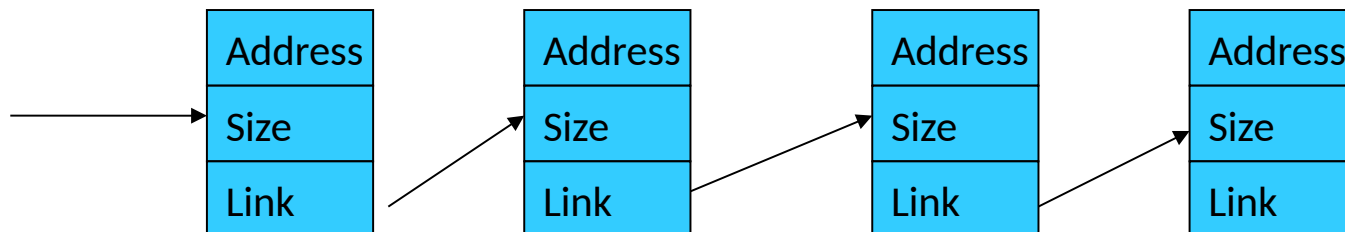


Four neighbor combinations for the terminating process X

Dynamic Partitioning: Placement

- First-fit algorithm

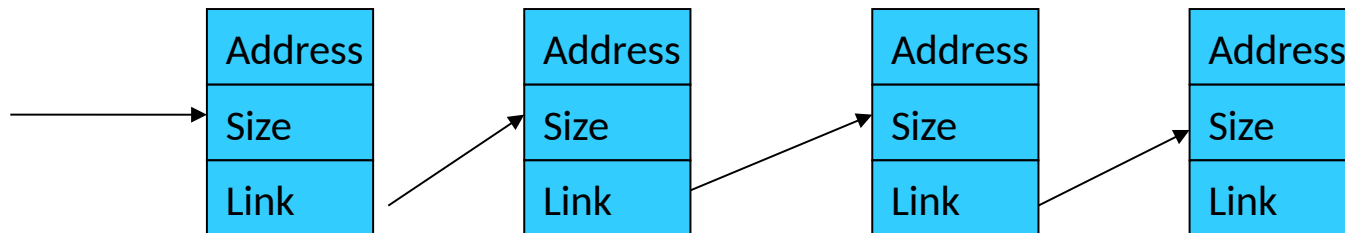
- Scan the list for the first entry that fits
 - If greater in size, break it into an allocated and free part
 - Intent: Minimise amount of searching performed
- Aims to find a match quickly
- Biases allocation to one end of memory
- Tends to preserve larger blocks at the end of memory



Dynamic Partitioning: Placement

- Next-fit

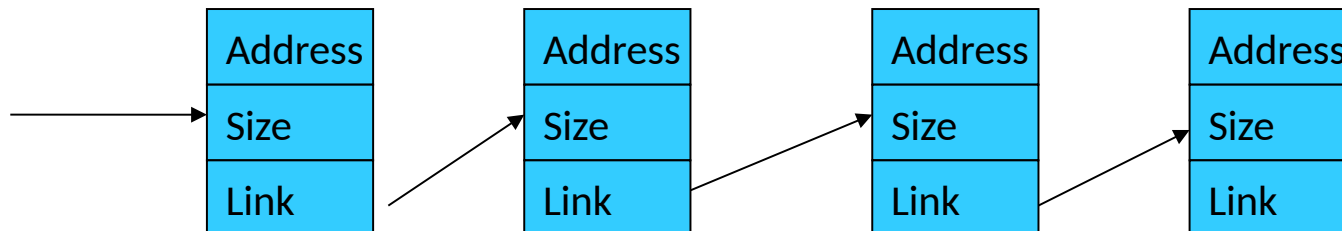
- Like first-fit, except it begins its search from the point in list where the last request succeeded instead of at the beginning.
- (Flawed) Intuition: spread allocation more uniformly over entire memory to avoid skipping over small holes at start of memory
- Performs worse than first-fit as it breaks up the large free space at end of memory.



Dynamic Partitioning: Placement

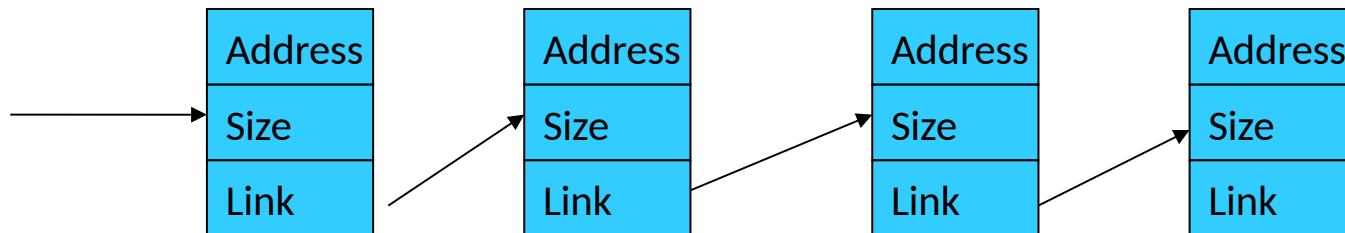
- Best-fit algorithm

- Chooses the block that is closest in size to the request
- Performs worse than first-fit
 - Has to search the complete list
 - CPU time wasted scanning the list.
 - Since the smallest block is chosen for a process, the smallest amount of external fragmentation is left.
 - Tends to create lots of small useless fragments.



Dynamic Partitioning: Placement

- Worst-fit algorithm
 - Chooses the block that is largest in size (worst-fit)
 - (whimsical) idea is to leave a usable fragment left over
 - Poor performer
 - Has to do more work (like best fit) to search complete list
 - Does not result in significantly less fragmentation



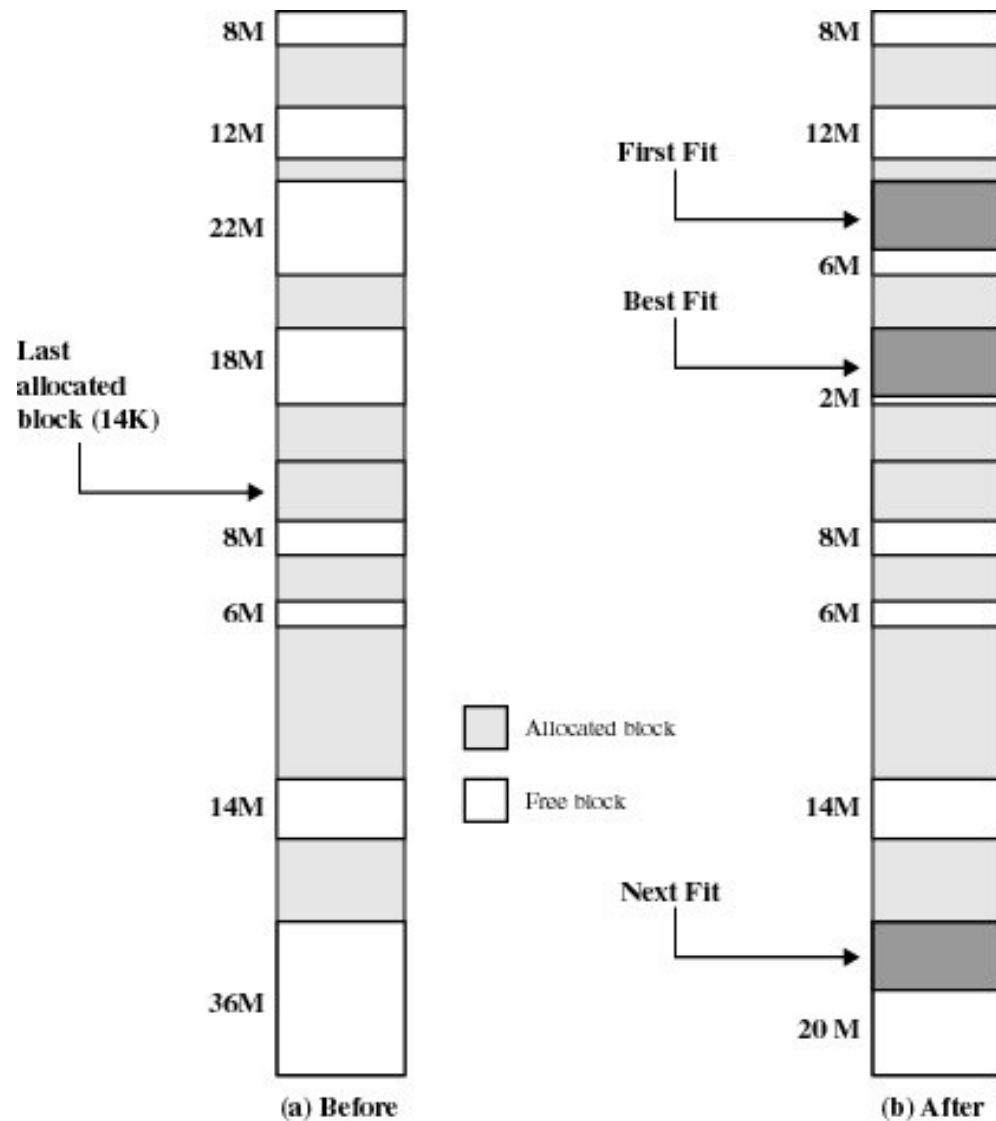


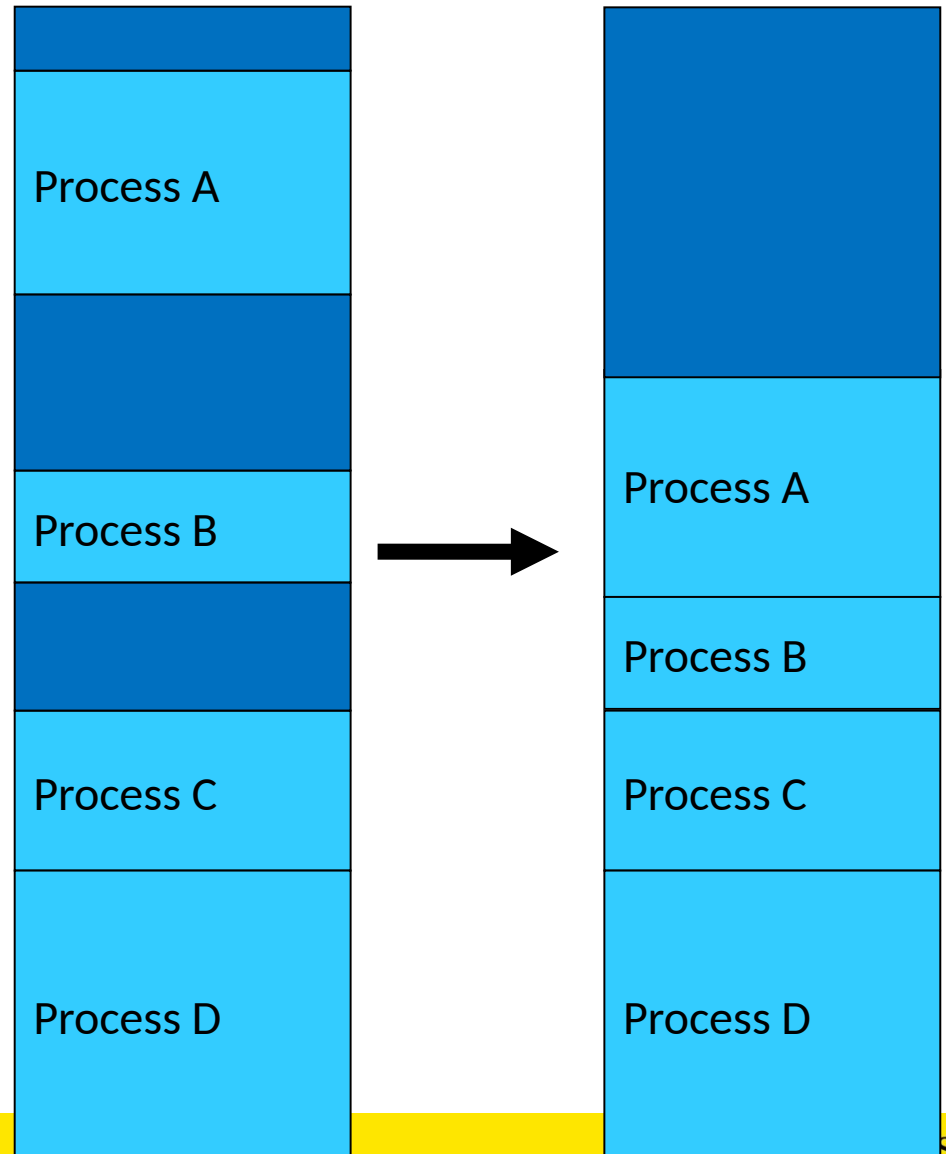
Figure 7.5 Example Memory Configuration Before and After Allocation of 16 Mbyte Block

Partition Allocation Algorithms

- Summary
 - First-fit generally better than the others and easiest to implement
- You should be aware of these
 - They are simple solutions to the problem of allocating memory, which is still a real problem.
 - Similar approaches can be used in disk block allocation.
- Note: These have been superseded by more complex and specific allocation strategies
 - Typical in-kernel allocators used are *lazy buddy*, and *slab* allocators.

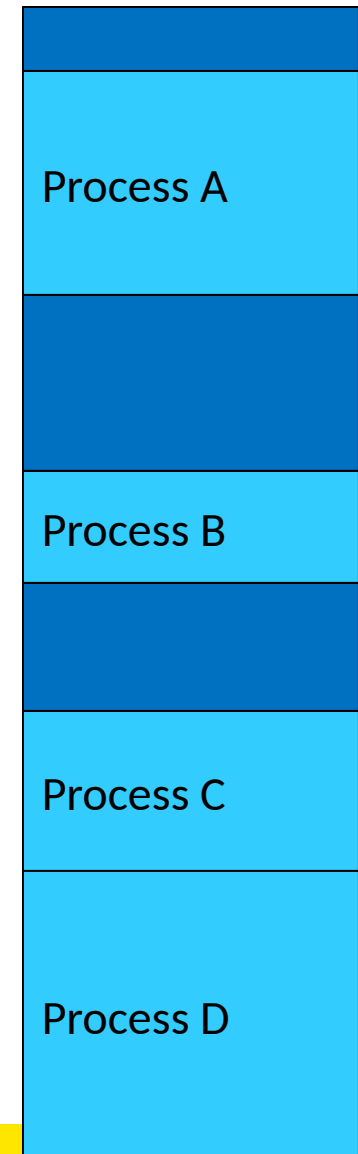
Compaction

- We can reduce external fragmentation by compaction.
 - Shuffle memory contents to place all free memory together in one large block.
 - Only if we can relocate running programs?
 - Pointers?
 - Generally requires hardware support.
- This motivates the more general approaches.



Some Remaining Issues with Dynamic Partitioning

- We have ignored
 - Relocation
 - How does a process run in different locations in memory?
 - Protection
 - How do we prevent processes interfering with each other?



Example Logical Address-Space Layout

- Logical addresses refer to specific locations within the program's memory
- Once running, these address must refer to real physical memory
- When are logical addresses bound to physical?

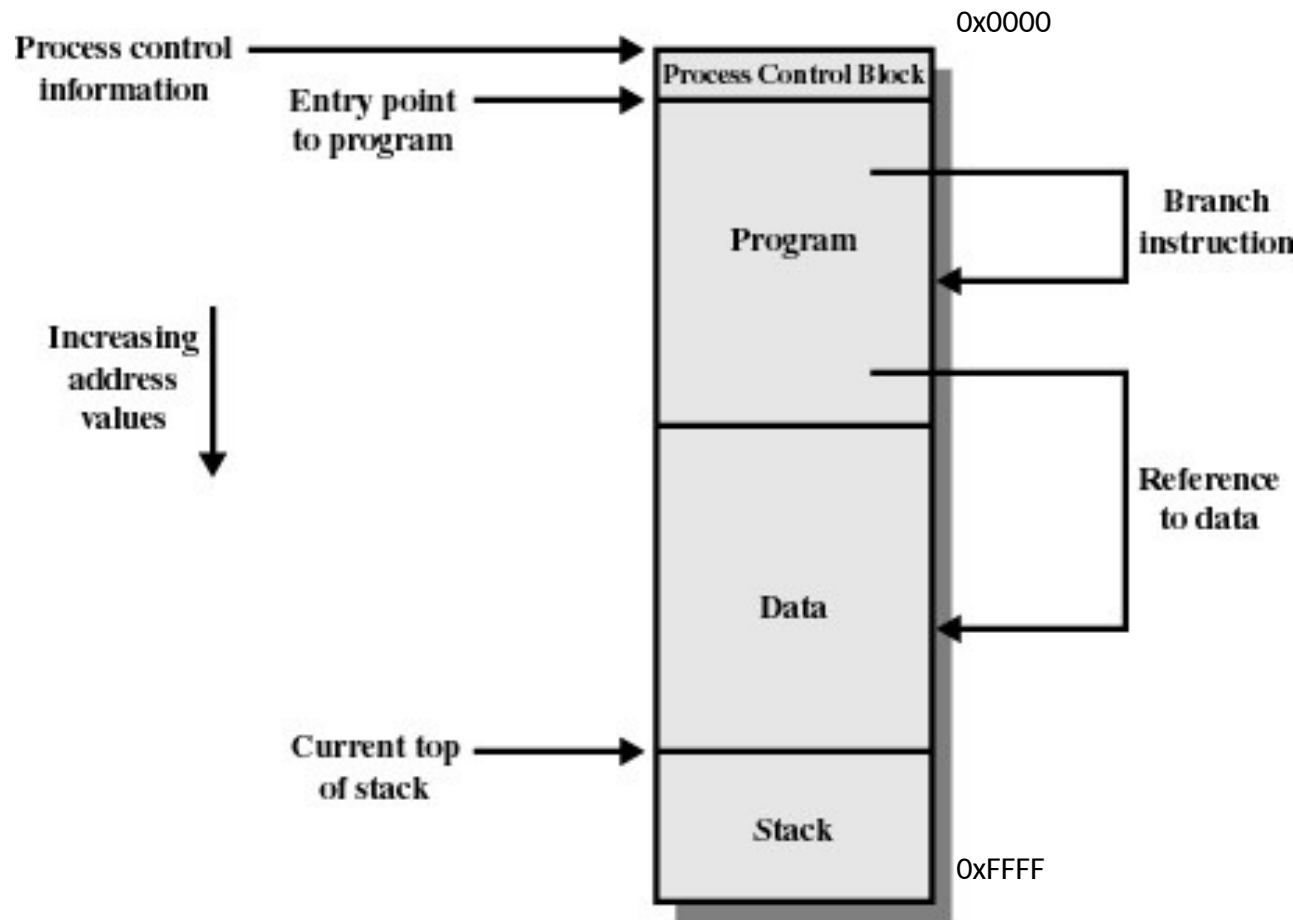
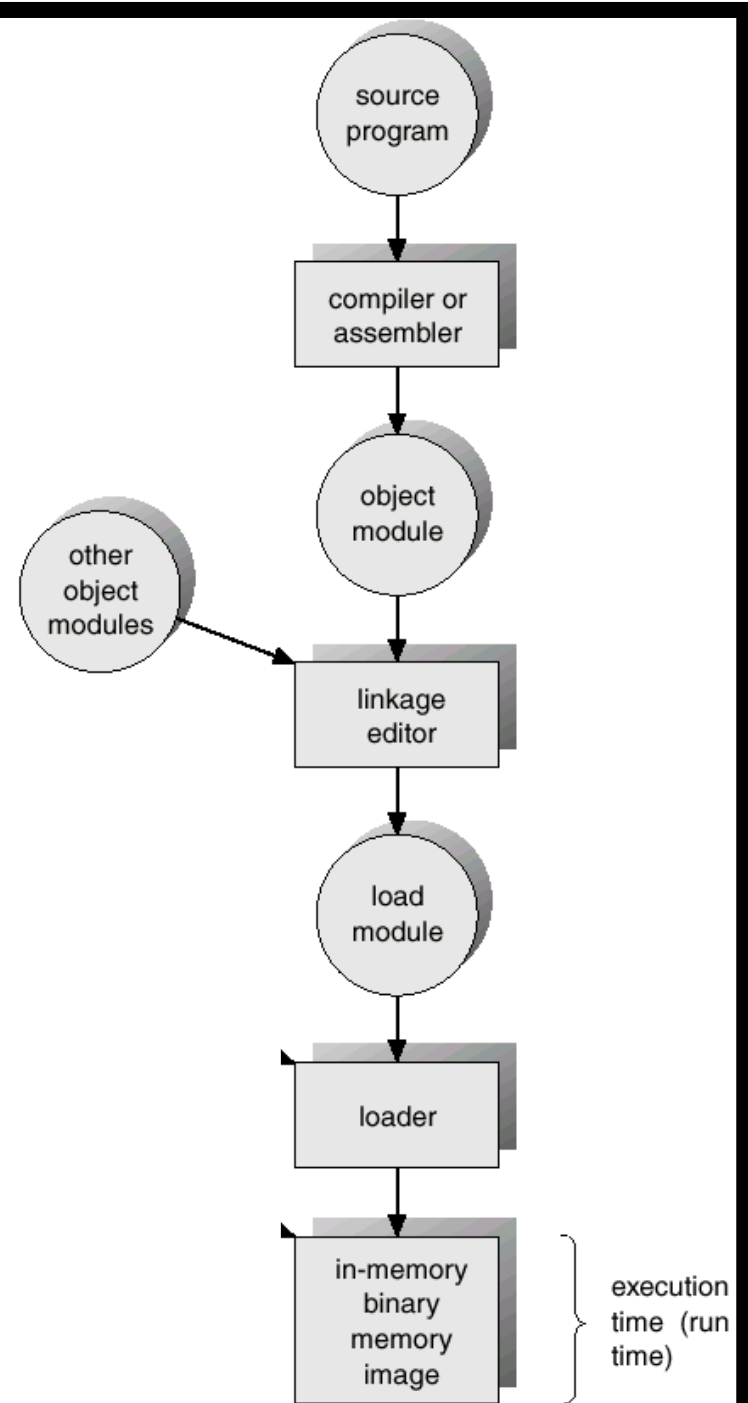


Figure 7.1 Addressing Requirements for a Process

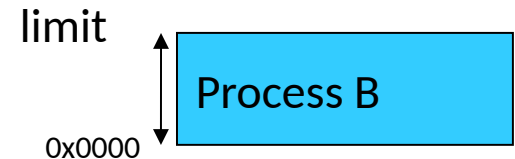
When are memory addresses bound?

- Compile/link time
 - Compiler/Linker binds the addresses
 - Must know “run” location at compile time
 - Recompile if location changes
- Load time
 - Compiler generates *relocatable* code
 - Loader binds the addresses at load time
- Run time
 - Logical compile-time addresses translated to physical addresses by *special hardware*.



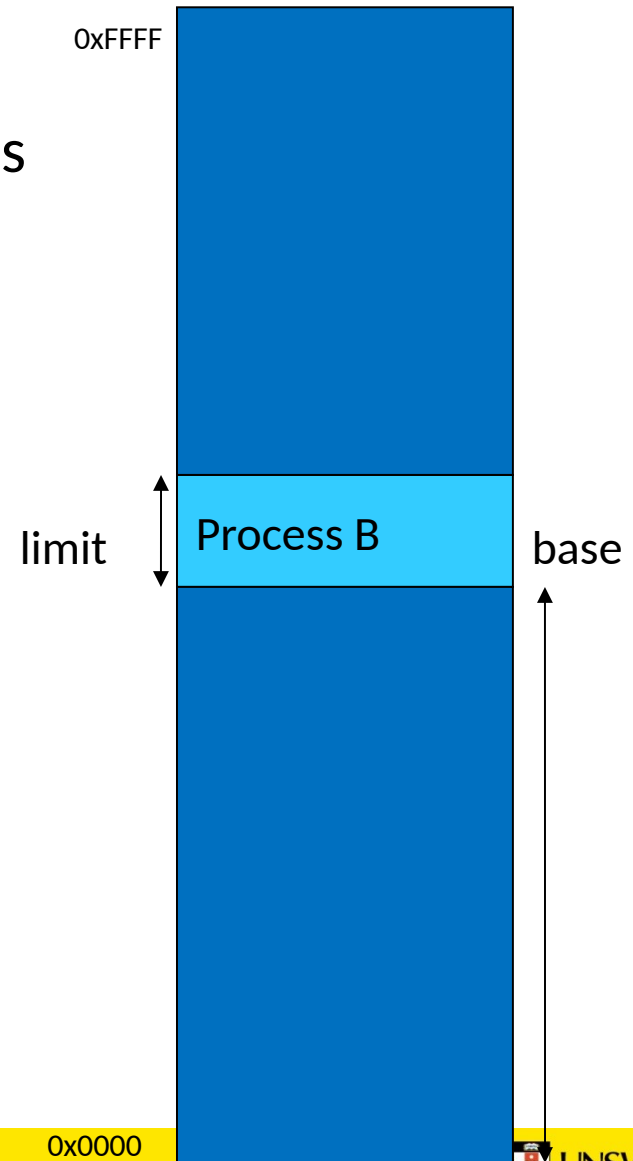
Hardware Support for Runtime Binding and Protection

- For process B to run using logical addresses
 - Process B expects to access addresses from zero to some limit of memory size

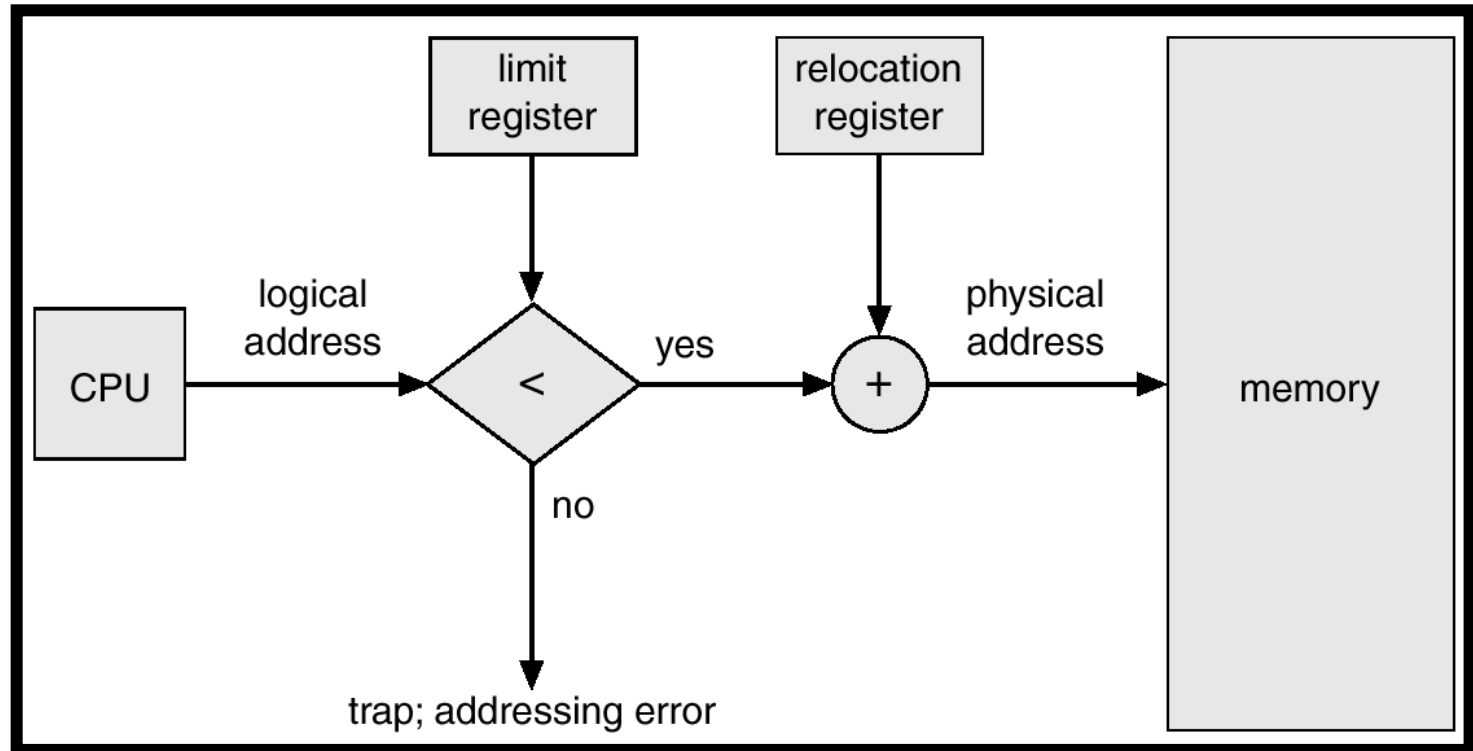


Hardware Support for Runtime Binding and Protection

- For process B to run using logical addresses
 - Need to add an appropriate offset to its logical addresses
 - Achieve relocation
 - Protect memory “lower” than B
 - Must limit the maximum logical address B can generate
 - Protect memory “higher” than B



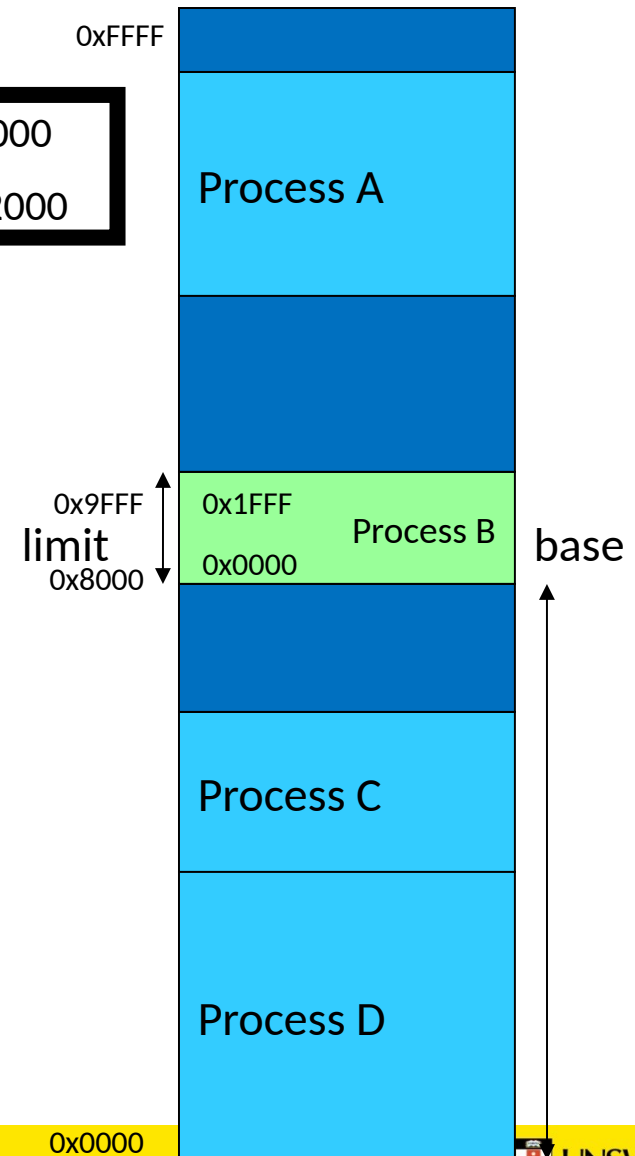
Hardware Support for Relocation and Limit Registers



Base and Limit Registers

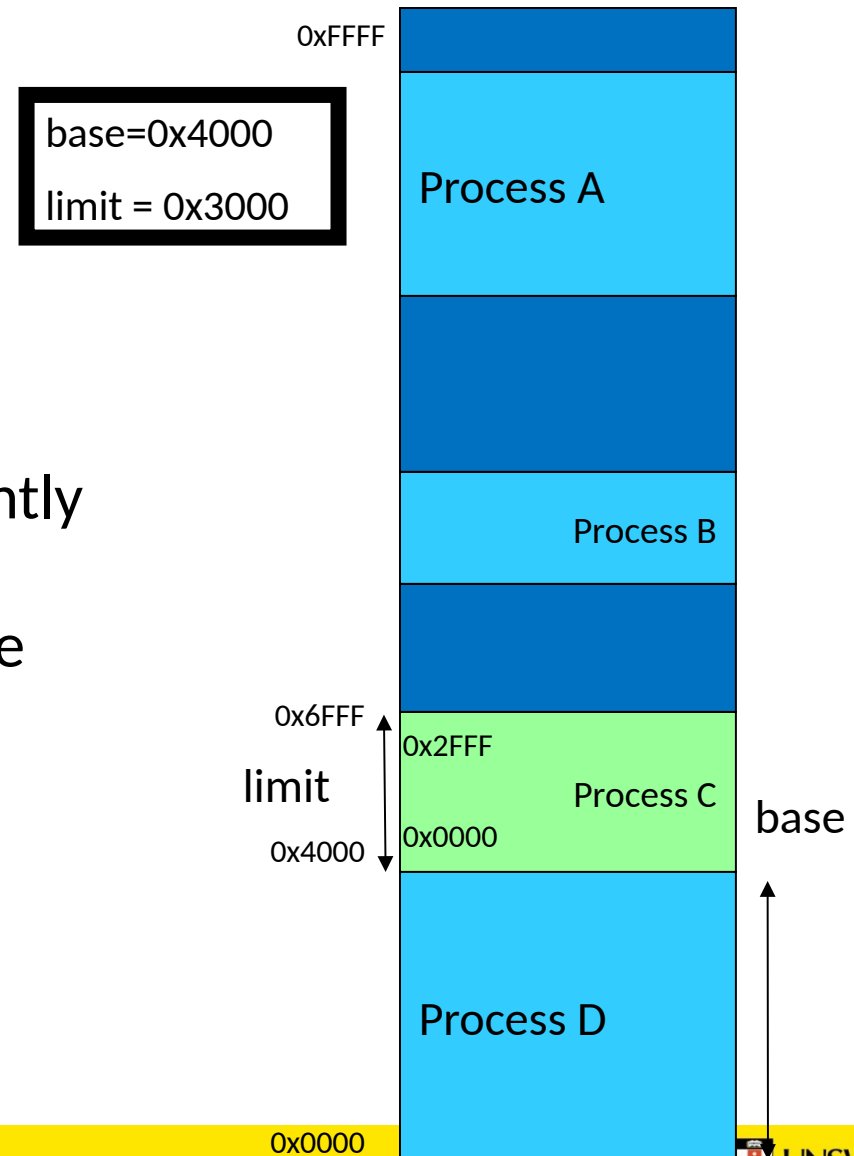
- Also called
 - Base and bound registers
 - Relocation and limit registers
- Base and limit registers
 - Restrict and relocate the currently active process
 - Base and limit registers must be changed at
 - Load time
 - Relocation (compaction time)
 - On a context switch

base=0x8000
limit = 0x2000



Base and Limit Registers

- Also called
 - Base and bound registers
 - Relocation and limit registers
- Base and limit registers
 - Restrict and relocate the currently active process
 - Base and limit registers must be changed at
 - Load time
 - Relocation (compaction time)
 - On a context switch

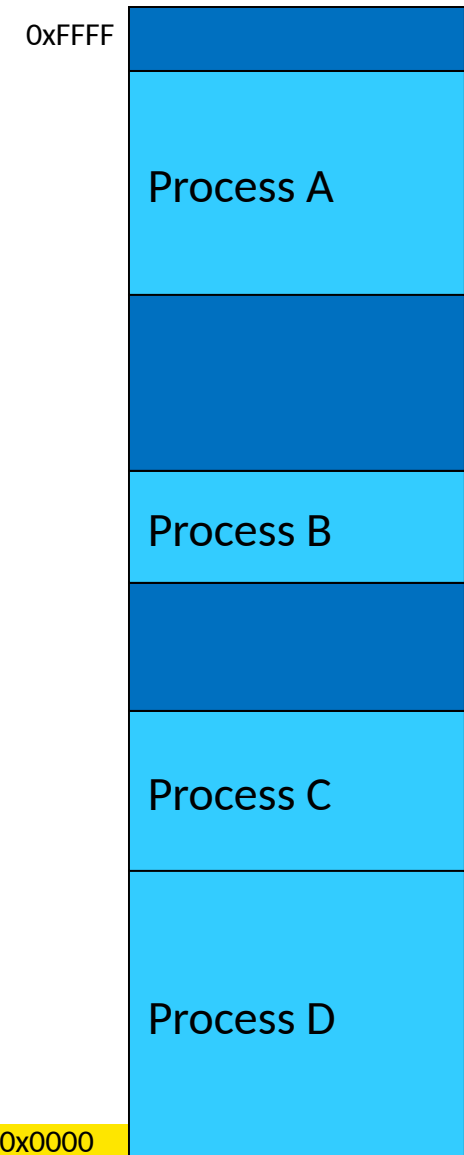


Base and Limit Registers

- Pros:
 - Supports protected multi-processing (-tasking)
- Cons:
 - Process memory must be contiguous.
 - Creates allocation/fragmentation issues.
 - Other features would be nice, e.g. sharing parts of address spaces.

Timesharing

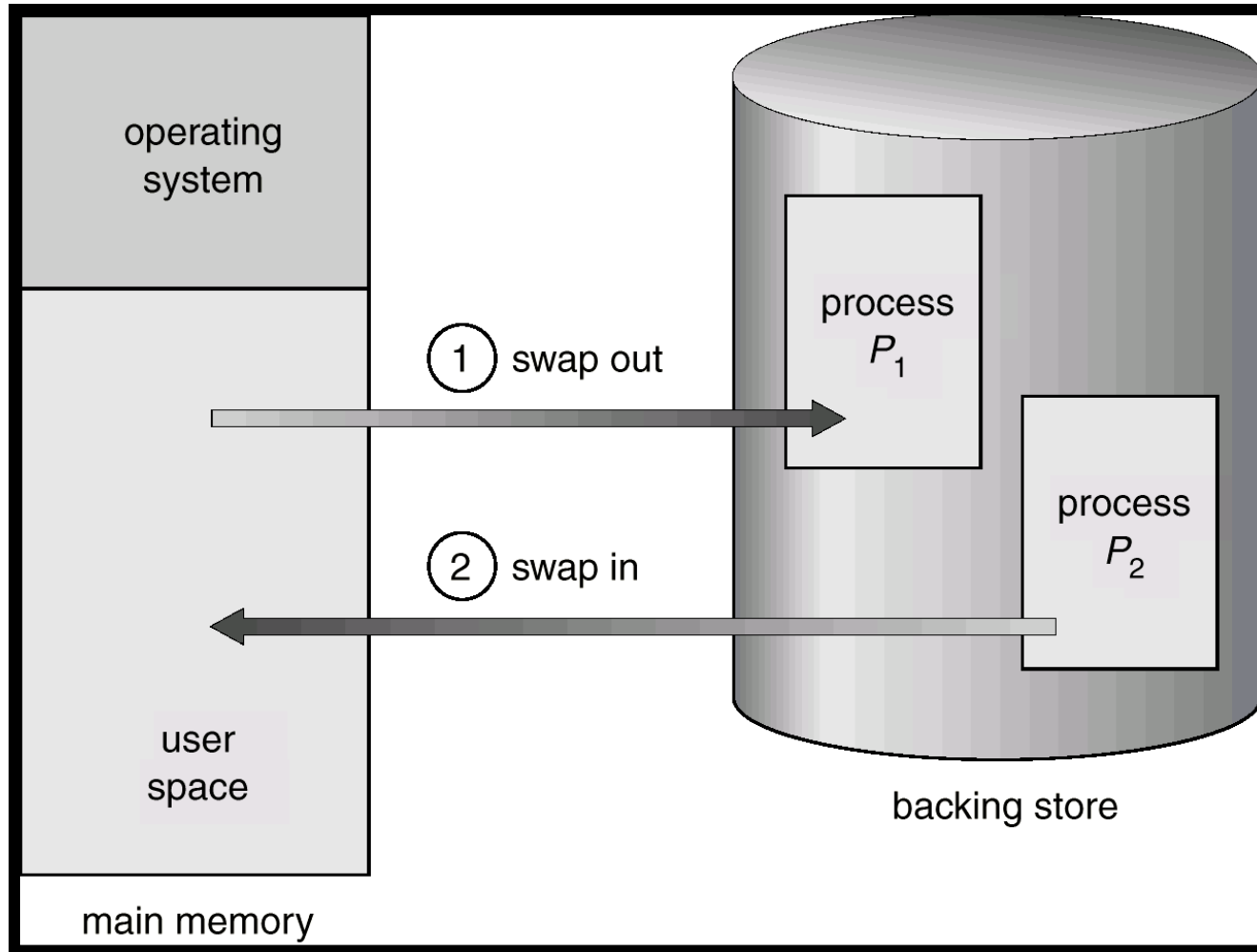
- Thus far, we have a system suitable for a batch system.
 - Limited number of dynamically allocated processes.
 - Enough to keep CPU utilised.
 - Relocated at runtime.
 - Protected from each other.
 - Processes finish and are cleaned up.
- Next step of the evolution: time-sharing.
 - A bit similar to modern cloud systems.
 - Some processes run periodically for a long period of time.
 - Some processes wait for various input.
 - Need to support a mix of active and inactive processes, of varying longevity



Swapping

- A process can be *swapped* temporarily out of memory to a *backing store*, and then brought back into memory for continued execution.
- Swapping involves transferring the whole process
- Backing store – fast disk large enough to accommodate copies of all memory images for all users; must provide direct access to these memory images.
- Can prioritize – lower-priority process is swapped out so higher-priority process can be loaded and executed.
- Major part of swap time is transfer time; total transfer time is directly proportional to the *amount* of memory swapped.
 - slow

Schematic View of Swapping



- So far we have assumed that every process *can* fit in memory.
- What can we do if a process is larger than main memory?

Virtual Memory

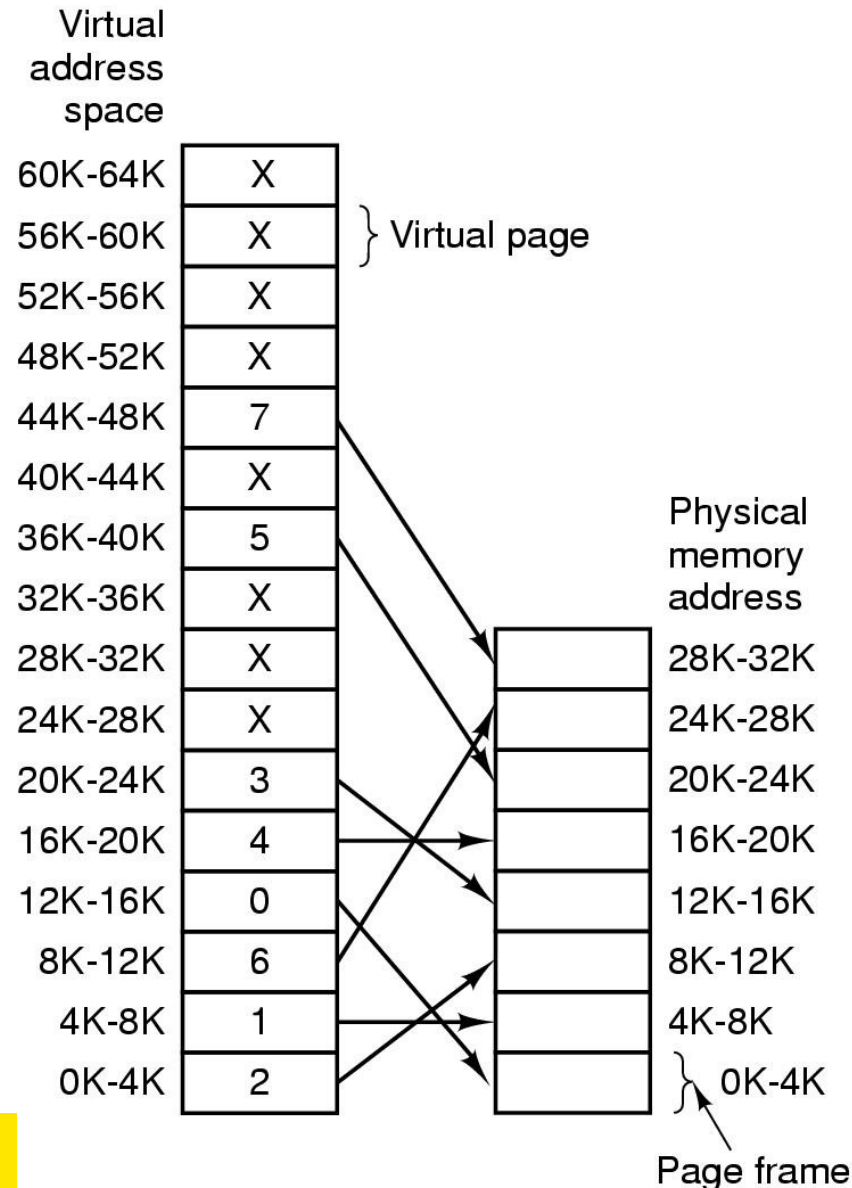
- Developed to address the issues identified with the simple schemes covered thus far.
- Like in an index-based file-system, permit the process's actual memory contents to be scattered across real memory.
- Creates some bonus features.
 - Fine-grained swapping to disk.
 - Partial sharing of memory.
- The process's view of memory is a virtual system.

Virtual Memory

- Two classic variants
 - Paging
 - Segmentation
 - (no longer covered in course, see textbook if interested)
- Paging is now the dominant one of the two
 - We'll focus on it
- Some architectures support hybrids of the two schemes
 - E.g. Intel IA-32 (32-bit x86)
 - Becoming less relevant

Virtual Memory – Paging Overview

- Divide physical memory into small chunks.
 - Called *frames*
 - Evenly sized
- Divide each process's virtual address space.
 - Called *pages*, same size as frames
 - Virtual memory addresses consist of a *page number* and *offset* within the page
- OS maintains a *page table*
 - contains the frame location for each page
 - Used by *hardware* to translate each virtual address to physical address
 - The mapping from virtual addresses and physical memory addresses is given by page table
- Process's physical memory does **not** have to be contiguous



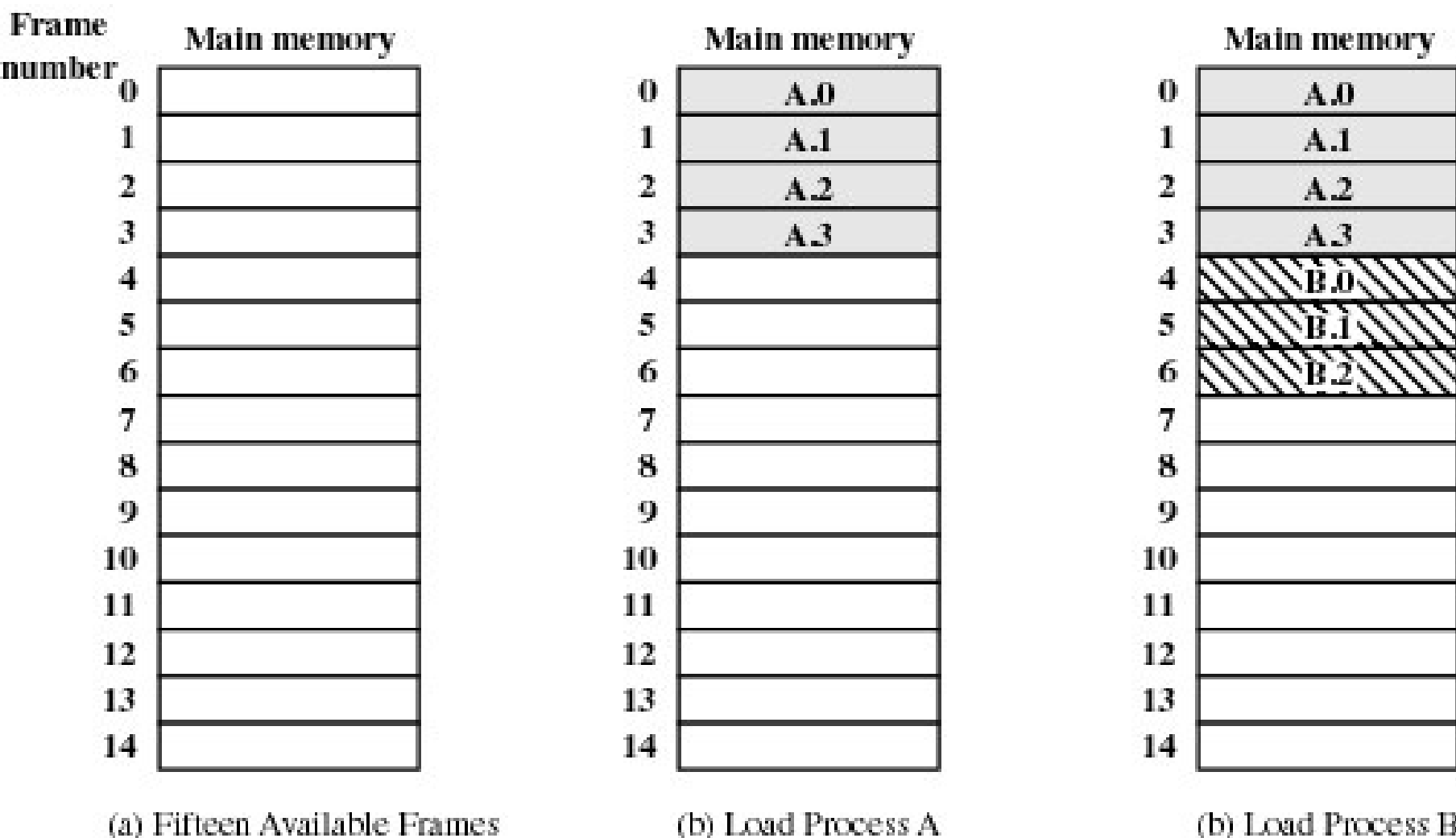


Figure 7.9 Assignment of Process Pages to Free Frames

Main memory	
0	A.0
1	A.1
2	A.2
3	A.3
4	B.0
5	B.1
6	B.2
7	C.0
8	C.1
9	C.2
10	C.3
11	
12	
13	
14	

(d) Load Process C

Main memory	
0	A.0
1	A.1
2	A.2
3	A.3
4	
5	
6	
7	C.0
8	C.1
9	C.2
10	C.3
11	
12	
13	
14	

(e) Swap out B

Main memory	
0	A.0
1	A.1
2	A.2
3	A.3
4	D.0
5	D.1
6	D.2
7	C.0
8	C.1
9	C.2
10	C.3
11	D.3
12	D.4
13	
14	

(f) Load Process D

0	0
1	1
2	2
3	3

Process A
page table

0	—
1	—
2	—

Process B
page table

0	7
1	8
2	9
3	10

Process C
page table

0	4
1	5
2	6
3	11
4	12

Process D
page table

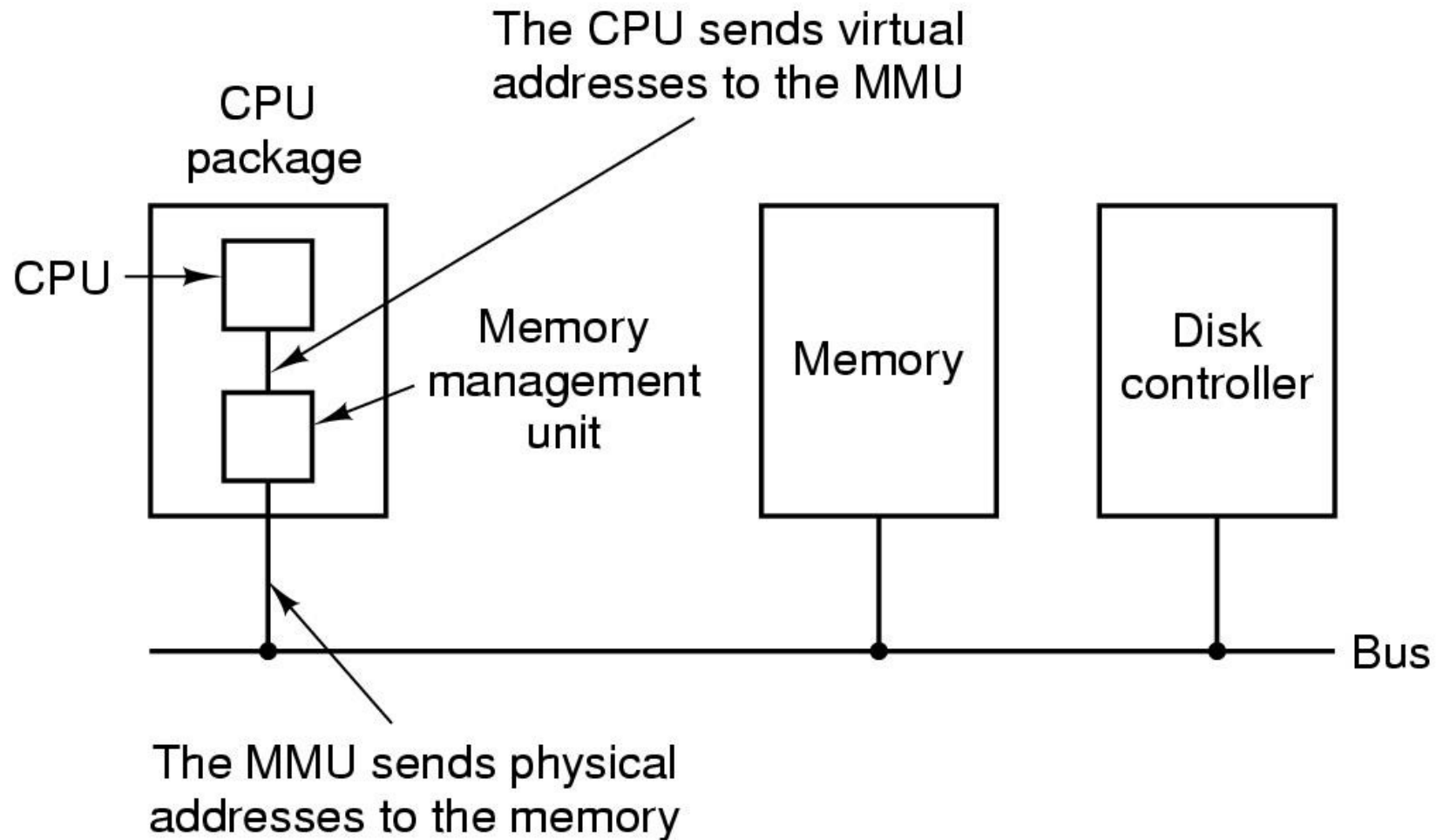
13
14

Free frame
list

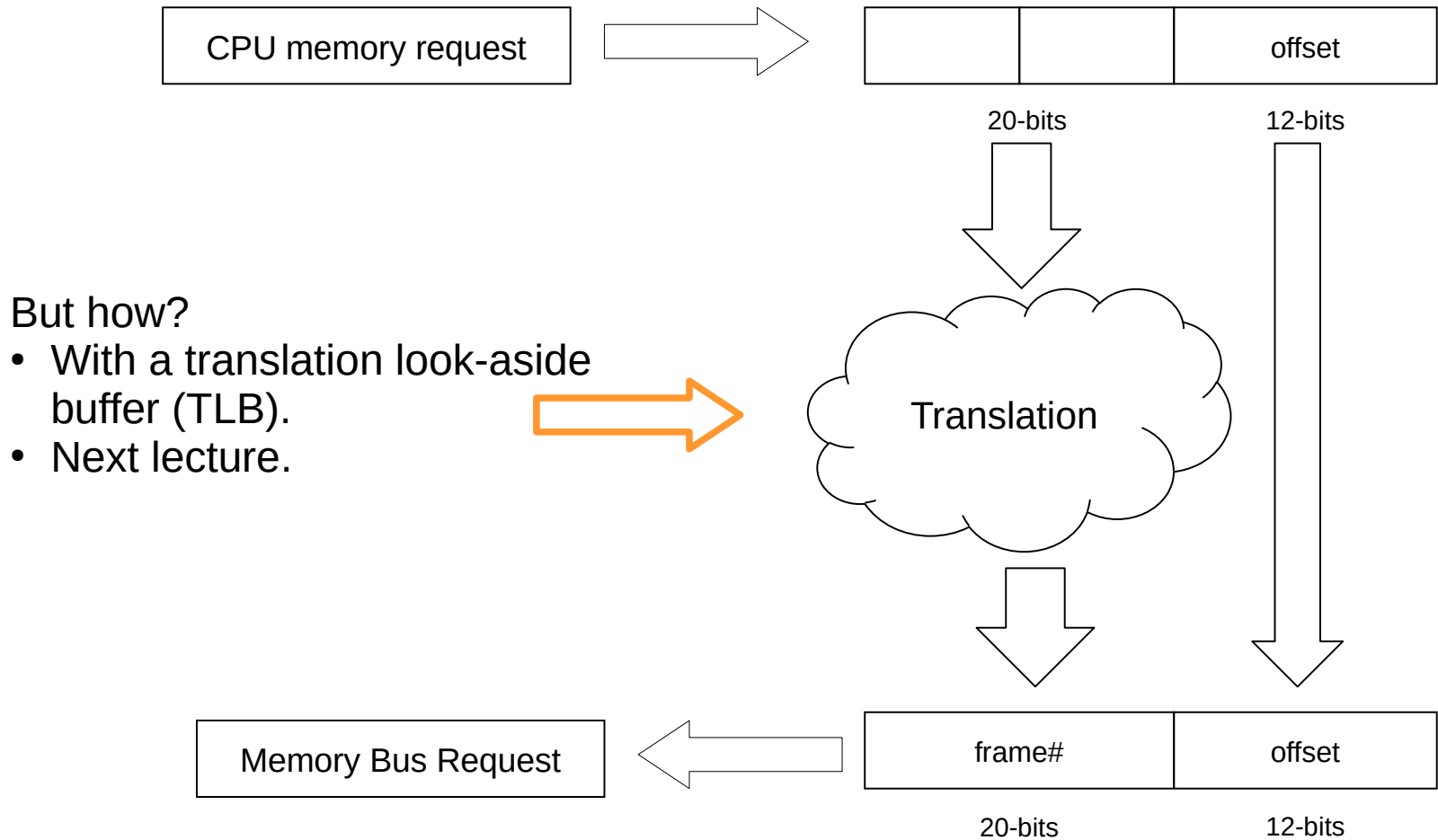
Paging

- No external fragmentation.
- Small amounts of internal fragmentation.
 - Rounding up process memory region to a page.
- Allows sharing by *mapping* several pages to the same frame.
- Abstracts the physical organisation of memory.
 - The programmer only deal with uniform virtual addresses.

Memory Management Unit (MMU)



MMU in Operation



Today's Lecture

(assuming we get to the finish)

- The need to manage system memory.
- More on allocation and fragmentation.
- Various schemes for static and dynamic allocation of regions.
- Simple relocation approaches.
 - Base and limit.
 - Offset.
- Paging.

Virtual Regions

(if we run over the end)

Paging:

- No need for the physical footprint of a process to be contiguous.
- No need for the range of available physical frames to be contiguous.
- No need for virtual address space to be contiguous. How might that be useful?

