

Topic Notes: File Systems

Persistent Storage

Recall that main memory is volatile – it loses its contents when the computer is shut off. It is also limited in size. Larger *persistent* storage devices are used for longer term storage and to hold data that is too large to fit entirely in main memory.

Persistent storage has taken many forms over the years:

- Paper: cards and tape with holes punched
- Magnetic tape
- Magnetic film on plastic (*floppy* disks)
- Magnetic film on metal (*hard* disks)
- Holes burned in metal film (CD/DVD *optical* disks)
- Quantum tunneling (*flash* memory)

We next look a bit more closely at the last three, which are the most common in use today.

Magnetic Hard Disks

Hard disks have been the standard persistent storage device in most computers for many years.

We'll concentrate on magnetic disks (floppy disk, hard disk). A hard disk may have multiple surfaces, or *platters*.

A read/write head (an *armature*) is needed for each platter. The platter has a magnetic film that holds the data. Each bit of data requires a magnetic particle that can be oriented by the read/write head when the bit is written, and whose orientation can be detected by the read/write head when the data is read.

The data on a disk is arranged in concentric rings called *cylinders* or *tracks*.

Each cylinder of the disk is divided into chunks called *sectors* that contain *blocks*, the minimum allocatable and addressable unit on the disk. Since there is more space on the outside of the disk, there may be more blocks in outer cylinders than there are on inner cylinders.

The particular configuration of cylinders, sectors and the number of platters is the *drive geometry*.

So to read or write data on the disk, a cylinder and sector must be specified. The read/write head must be positioned over the desired cylinder and sector. The read/write heads are typically connected to the end of a moveable arm. This arm is moved to position the head at the correct cylinder. When the disk rotates and the desired sector reaches the read-write head, the read or write operation can proceed.

The speed of this operation depends on two major factors:

- *seek time* – the time it takes to move the read/write head to the correct cylinder
- *rotational latency* – the time it takes for the correct sector to rotate under the read/write head

Typical platter rotational speeds are 5400, 7200, but some can be as high as 15,000 RPM.

Access times are measured in milliseconds. Given that CPU operations are measured in picoseconds, disk access is incredibly slow by comparison.

We can minimize seek time by minimizing the distance the read/write head has to move in order to service the incoming requests.

Typical capacities are now measured in hundreds of gigabytes to several terabytes.

Optical Data Storage

Compact Disc (CD) and Digital Video Disc (DVD) storage is also very common in modern computers.

An optical storage media like a CD or a DVD is simply a rotating plastic disc with metal film on it.

Unlike hard disks, where the data is arranged in concentric cylinders on the disk, optical storage organizes the data in a continuous spiral (like a vinyl record).

The bits in this case are stored as *flats* (or *lands*) and *holes* (or *pits*) which are burned into the metal film. A laser (780 nm wavelength for a CD) on an armature bounces off of these lands and pits to read the bits.

A CD has a data capacity of 650 to 900 MB.

A DVD uses a 650 nm laser and a Blu Ray a 405 nm laser, which means the pits can be smaller. Data capacities for a DVD are 8+ GB. A dual-layer Blu Ray data disc can hold 50 GB of data.

Access times are in the range of hundreds of milliseconds.

Flash Storage

Flash storage is a non-volatile memory typically found in some (typically smaller capacity) external drives and memory cards, but is sometimes used for internal storage in *solid state* drives.

This type of storage uses a *quantum tunneling* to “trap” an electronic bit of information between two insulators. This bit retains its value even when the device is not powered.

Capacities of flash storage now range from around 64 MB to 256 GB. Access times are typically in the 1 to 10 millisecond range.

An important long-term consideration when using flash memory is that it can “wear out” from repeated writing. Many flash devices are guaranteed to withstand 100,000 write cycles, some up to one million.

Further details of flash storage are not our concern here.

File System Interface

Our next area of focus will be on how to organize information on disks. Some of this organization is independent of the underlying technology, but some is specific.

Hopefully everyone has a good idea what we mean by a *file*.

- Files can be data or programs.
- Files can be simple or complex (plain text, or a specially-formatted file).
- The structure of a file is determined by a combination of the OS and the program that creates it.
- Files are stored in a *file system*, which may exist
 - on a disk,
 - on a tape,
 - in main memory,
 - or, any other storage device that might be available.

Files have a number of attributes:

- filename
- file type (maybe)
- location – where is it on the device
- size
- protection/permissions (maybe)
- timestamp, ownership
- directory information

A number of operations can be performed on files, many of which you have have been doing throughout your lives as computer users without giving it much thought. As programmers, you have thought somewhat more carefully than the average user.

- create
- write/append
- read
- seek (reposition within file)
- delete
- truncate
- open, close

File Types

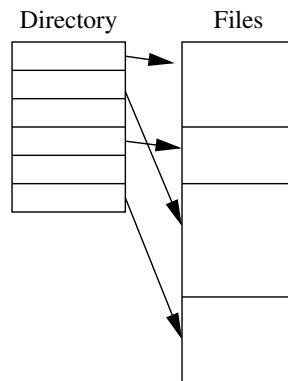
How does a “type” get assigned to a file?

- One can use a file extension (`.c`, `.exe`, `.doc`, `.tex`, `.mp3`, etc.)
- However, a file extension may or may not be functionally important – even if not, the extension can give the OS and the user an hint as to the type of the file
 - Windows uses an unenforced file extension registration
 - Macintosh can enforce types within a file – using a special part of a file called a “resource fork” to store extra information including the application that created it
 - Unix uses “magic numbers”. See the `file` command and `/usr/shared/misc/magic` in FreeBSD, `/etc/magic` in Solaris, files within `/usr/share/file/magic` on Mac OS X.

File access may be *sequential* or *direct*. Tapes support only sequential access, disk files may support both.

Directories

A listing of the files on a disk is a *directory*.



Both the directory structure and the files reside on the disk.

The directory may store some or all of the file attributes we discussed.

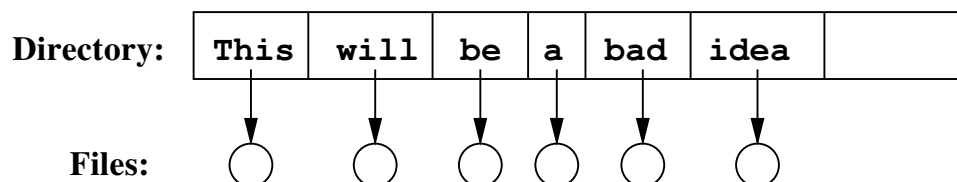
A directory should be able to support a number of common operations:

- search for a file
- create a file
- delete a file
- rename a file
- listing of files
- filesystem traversal (`cd`)

There are many ways to organize a directory, with different levels of complexity, flexibility, and efficiency. We will look at several possibilities.

- Single-Level Directory

The simplest method is to have one big list of all files on a disk.



This can be used for a simple system. A disk for the C-64 worked like this:

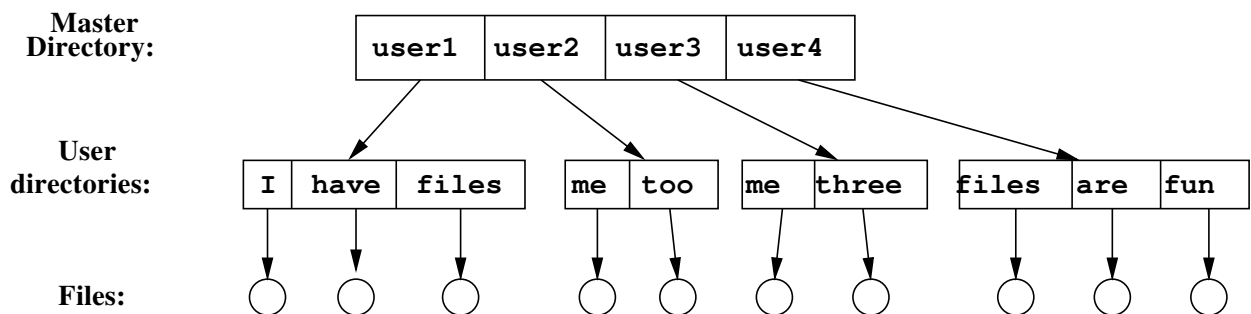
```
LOAD "$", 8
LIST
```

This command reads the special file \$, which indicates the disk's directory, into memory (like a BASIC program) from device 8 (the floppy drive). Then LIST lists the "program" in memory, which in this case is just a list of the files on the disk.

Of course, this model breaks down pretty quickly:

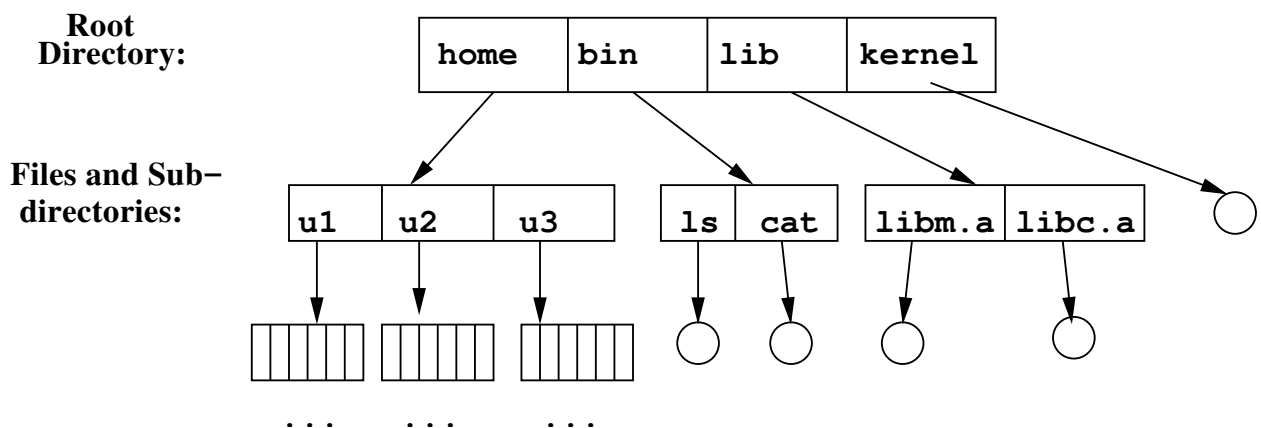
- There cannot exist two files with the same name, which would likely be necessary for multiple users/programs on a disk.
 - There is no way to group files – it's just one big list.
 - Searches for a file with a given name need to look through the entire directory.
- Two-Level Directory

We can create a separate directory for each user:



- Files now have a path name, *e.g.*, /user1/have.
 - Different users can have files of the same name (*e.g.*, /user2/me and /user3/me) distinguished by the path.
 - Searching is more efficient, as only one user's list needs to be searched.
 - But... still no grouping capability for an individual user's files.
- Tree-structured Directory

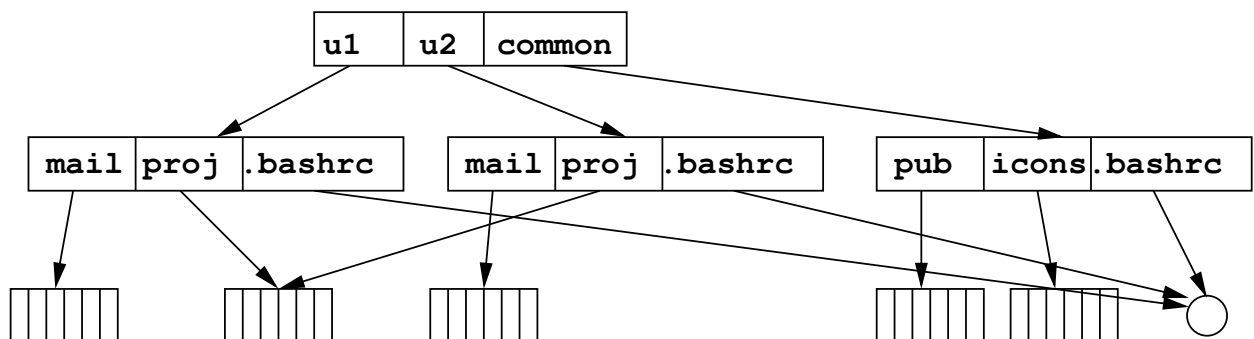
Now, we get to something more reasonable and usable:



- Any directory entry can be either a file or a *subdirectory*.
 - Files can be grouped appropriately.
 - Search is more efficient – follow the path to the file sought, and look only in that directory.
 - Here, we need to add the concept of a *current working directory* and related support:
 - * traverse directory (`cd`)
 - * operate on files in current directory by default
 - * or specify path, either *relative* or *absolute*
 - This scheme requires support for creation and removal of directories as well as files.
 - And new issues arise: *e.g.*, what happens when a non-empty directory is deleted?
- Acyclic-Graph Directories

The tree model does not allow the same file to exist in more than one directory. We can provide this by making the directory structure an acyclic graph.

Two or more directory entries can point to the same subdirectory or file, but in this model we still disallow any directory entry pointing “back up” the directory structure.



These kinds of directory graphs can be made using *links* in the Unix world, *shortcuts* in the Windows world, or *aliases* in the Mac world.

This allows multiple names for and multiple paths to the same file on the disk.

Unix links can be

- *symbolic* or *soft link* – specify a path to the file (logical) – `ln -s` – original file is “real” others are just pointing to that one
- *hard link* – actual link to the same file on the disk from multiple directories (physical) – `ln` – all hard links are equal

This allows sharing of files, but introduces complications – what happens when the file is removed from one of the directories? If there may be more references to the file, can we delete it? With symbolic links, the file just gets deleted and we have a *dangling pointer*.

With hard links, a reference count is maintained, and the actual file is only deleted when all references to it are removed.

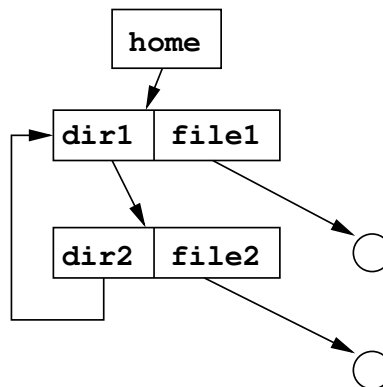
In-class demo: links

- General Graph Directory

What if we allow links back up the chain?

Unix directories have this built in – all directories except the system’s root directory have a special entry `..` that indicates the parent directory, and an entry `.` that indicates the current directory.

But they also allow links to be created back “up the chain” of the directory structures, potentially introducing cycles in the directory graph.



In-class demo: dirs

When general graph directories are allowed, we need to be careful with commands like `find` that search a directory and its subdirectories for something. The search is infinite if cycles are followed. Typically, a program like `find` will not follow symlinks.

Problematic cycles can be avoided by allowing “up” links to files, not directories. Could also run cycle detection every time a new link is added, if this is a concern. Unix leaves it up to programs to make sure they treat symlinks appropriately.

BSD 4.3 limits the number of links allowed to be traversed for any given path name to 8 to avoid undetected cycles. This limit is actually 32 in the current version of FreeBSD, and 20 for Solaris.

In the <https://github.com/SienaCSISOperatingSystems/morelinks-example> repository, we see how symlinks can be created with system calls, and from that we can create lots of links and test this limit.

Directory Implementation

An individual subdirectory will typically contain a list of files. How best to store this list?

- Linear list – a simple list of names, each of which has a pointer to the file’s data blocks. This is straightforward, but requires a costly search on large directories.
- Hash Table – hashed linear list – decreases search time, but more complex to implement.

Another consideration: *case sensitivity* of filenames. Modern Windows, MacOS filesystems have filenames that remember case, but searches are case insensitive (sometimes called *case preserving*). Most Unix filesystems are truly case sensitive.

Disks and Partitions

A system may have a number of disks, each with one or more *partitions*. These are logical subdivisions of the physical disk, often created to help better organize data on the disk.

Demo: `df -kl`

A partition is where a filesystem gets created (more on that soon). Once we have a filesystem, we need to make it accessible to the world.

In DOS/Windows, this typically involves assigning a letter to each partition. Then there is a directory hierarchy within each partition.

In Unix, there are no drive letters, everything is considered to be part of one big hierarchy. One partition forms the root directory (/) and all others are *mounted* into the structure it defines.

The places where partitions get mounted are called *mount points*, and are nothing more than regular directories. When a partition is mounted onto a mount point, the directory is replaced by the contents of the partition mounted.

Demo: mounting

When the mount point directory is accessed, the *virtual file system* layer of the OS notices that the directory has been used as a mount point, and sets the current directory to the root of the partition mounted there.

The partitions mounted can be of any type, but all appear to be part of the same directory structure. The system delivers requests to the appropriate partition, and a filesystem-type-specific set of operations are used to access the actual filesystem.

The list of partitions and their mount points and types are listed in a *file system table* file. In FreeBSD, it is located in `/etc/fstab`; in Solaris, it is `/etc/vfstab`.

The partitions mounted may be remote as well as local – more on this later.

Disk Partitioning

Why might we partition a disk?

- logical separation of types of files (bootable OS, system programs, home directory space,

shared space, scratch space) for security or backup purposes.

- want to run multiple OSs on the same system.
- separate partition to use as virtual memory (“swap partition”).
- to get around OS limits on the size of a filesystem when a single disk is larger than that limit.

How can we define these partitions? These are usually specified with a system disk management utility.

- DOS/Windows fdisk
- FreeBSD disklabel/bsdlabel
- MacOS Disk Utility

All of these do the same basic things. Break up the disk, usually on cylinder boundaries, into logical subunits.

Each of the partitions gets a device name, and in each of these we create a filesystem.

The filesystem can then be mounted at a given mount point (in the Unix world) or at a drive letter in DOS/Windows.

Demo: fdisk

File System Implementation

Suppose we have partitioned a disk, and it's time to take those disk blocks that have been reserved for our partition and create an actual *file system* to hold our files.

The OS could just provide access to the blocks and let programmers deal with everything, but that's not very nice.

We want to provide those things that operating systems are supposed to provide:

- convenience
- protection
- efficiency

Several issues need to be considered:

- how do we allocate disk blocks within our partitions to files and directories?
- how we decide what blocks are available?

- what is the complexity and efficiency of the choices we make for those?

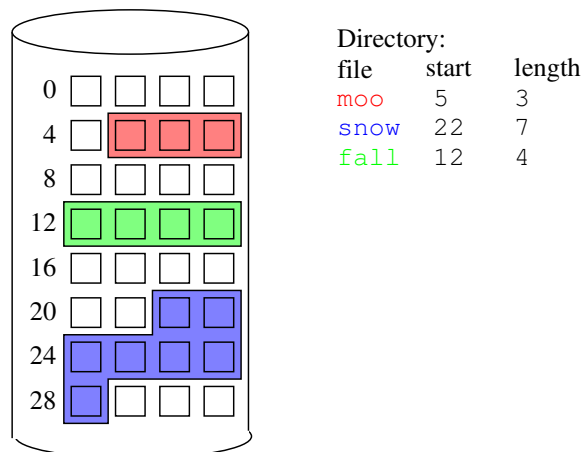
We have already talked about possible directory structures. Most likely, these directory structures will be implemented as files at some level. We'll need to be able to find them on the disk just like other files.

Allocation Methods

So first, we consider how disk blocks are allocated to files.

- Contiguous Allocation

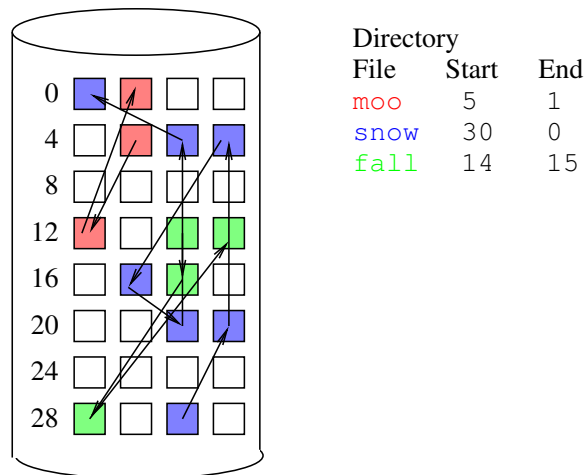
Each file is allocated a set of contiguous disk blocks



- similar to contiguous allocation of memory
- simple – directory entry needs only starting location (block number) and length (number of blocks)
- supports random access into files – can easily compute and read the block that contains a certain part of the file.
- can lead to holes (external fragmentation)
- may be difficult to have a file grow
- reading should be very efficient, since consecutive blocks of the file can be stored in consecutive blocks on the disk
- Clustered allocation, or Extents

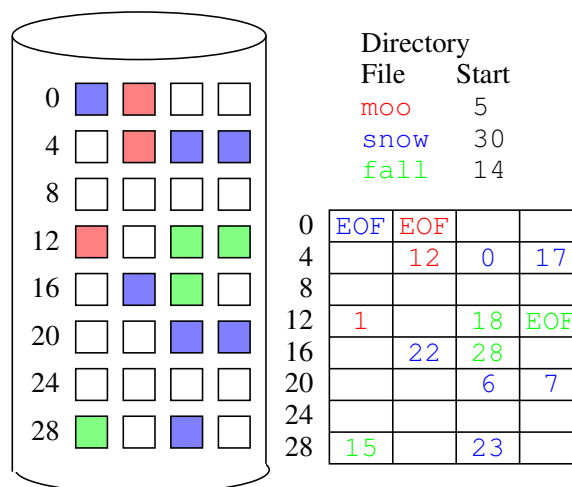
This approach is analogous to segmentation for memory allocation. Files are allocated as a collection of clustered blocks, or *extents*, which are contiguous chunks of disk blocks. Each has a starting block and a size.
- Linked Allocation

Each disk block has a pointer to the next disk block in the file as well as some file data.



- need to reserve part of each data block for a pointer – can make for odd-sized data blocks
- directory entry requires only starting block
- easy to append to a file
- no external fragmentation
- no random access – have to traverse each block
- a bad disk block means the entire file from that block on is lost

A variation on this is the *File Allocation Table (FAT)* used by MS-DOS and pre-NT Windows versions. This gathers the links into one table.

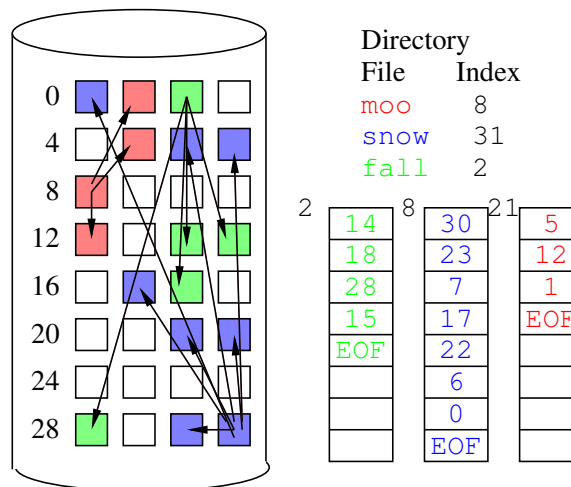


- get to use the whole disk block for data
- a bad disk block means only that block is lost
- unless... the FAT itself goes bad, in which case we have a problem – have backup copies on the disk, then run your favorite rescue program

- somewhat better random access – traverse the FAT only – read disk blocks only for the data stored there
- each disk block needs a FAT entry – total number of blocks, in turn total size of a partition – is limited by the size of the FAT
- increased block size means fewer blocks/FAT entries, but more internal fragmentation

- Indexed Allocation

Use disk blocks as *index blocks* that don't hold file data, but hold pointers to the disk blocks that hold file data.



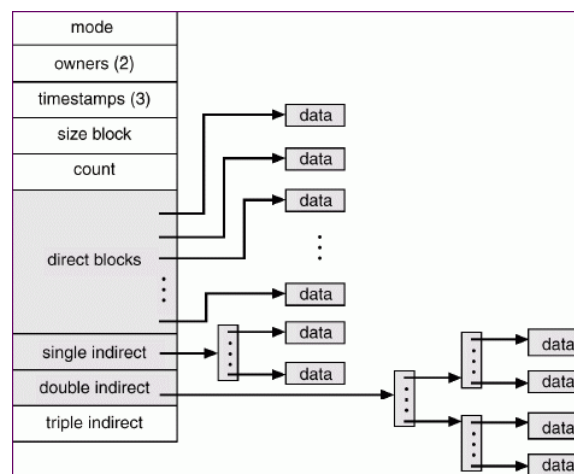
- directory entry now contains a pointer to the index block
- each file's index block contains pointers to all of its data blocks
- random access is similar to FAT
- a bad data block costs only that block, bad index block could cost the entire file
- size of a file is limited by the number of pointers a data block can hold – if a block holds 512 bytes, and a pointer to a disk block takes 2 bytes, we are limited to 256-block, or 128 KB files
- now even small files require two data blocks – extra disk reads, and potentially wasted space

Can get around the file size limitation in a few ways:

- *linked indexed allocation* – use the last entry in the index block as a pointer to another index block
 - * this removes file size limitations
 - * random access becomes a bit harder
- *two-level index* – the index block points only to other index blocks

- * file size limitation is not as severe – for example above, disk file now are addressed by a 256-entry index block, each of which points to a 256-entry index block, meaning we can store 65536-block or 32 MB files.
- * random access is better
- * but all files take at least 3 blocks of space and access time
- Can add more levels for larger files
- Unix Inodes

Many Unix filesystems (Berkeley Fast Filesystem, Linux ext2fs, Sun ufs, ...) take an approach that combines some of the ideas above.



- each file is indexed by an *inode*
- inodes are special disk blocks set aside just for this purpose (see `df -i` to see how many of these exist on your favorite Unix filesystem)
- they are created when the filesystem is created
- the number of inodes limits the total number of files/directories that can be stored in the filesystem
- the inode itself consists of
 - * administrative information (permissions, timestamps, etc.)
 - * a number of direct blocks (typically 12) that contain pointers to the first 12 blocks of the file
 - * a single indirect pointer that points to a disk block which in turn is used as an index block, if the file is too big to be indexed entirely by the direct blocks
 - * a double indirect pointer that points to a disk block which is a collection of pointers to disk blocks which are index blocks, used if the file is too big to be indexed by the direct and single indirect blocks
 - * a triple indirect pointer that points to an index block of index blocks of index blocks...

- interesting reading on your favorite FreeBSD system: `/sys/ufs/ufs/dinode.h`
- small files need only the direct blocks, so there is little waste in space or extra disk reads in those cases
- medium sized files may use indirect blocks
- only large files make use of (and incur the overhead of) the double or triple indirect blocks, and that is reasonable since those files are large anyway
- since the disk is now broken into two different types of blocks – inodes and data blocks, there must be some way to determine where the inodes are, and to keep track of free inodes and disk blocks. This is done by a *superblock*, located at a fixed position in the filesystem. The superblock is usually replicated on the disk to avoid catastrophic failure in case of corruption of the main superblock

Disk Allocation Considerations:

- limitations on file size, total partition size
- internal, external fragmentation
- overhead to store and access index blocks
- layout of files, inodes, directories, etc, as they affect performance – disk head movement, rotational latency – many unix filesystems keep clusters of inodes at a variety of locations throughout the file system, to allow inodes and the disk blocks they reference to be close together
- may want to reorganize files occasionally to improve layout (disk defragmenting, etc)

Free Space Management

With any of these methods of allocation, we need some way to keep track of free disk blocks.

Two main options:

1. *bit vector* – keep a vector, one bit per disk block
 - 0 means the corresponding block is free, 1 means it is in use
 - search for a free block requires search for the first 0 bit, can be efficient given hardware support
 - vector is too big to keep in main memory, so it must be on disk, which makes traversal slow
 - with block size 2^{12} or 4KB, disk size 2^{40} or 1 TB, we need 2^{28} bits (16 MB) for the bit vector (seems reasonable)
 - easy to allocate contiguous space for files

2. *free list* – keep a linked list of free blocks

- with linked allocation, can just use existing links to form a free list
- with FAT, use FAT entries for unallocated blocks to store free list
- no wasted space
- can be difficult to allocate contiguous blocks
- allocate from head of list, deallocated blocks added to tail, both $O(1)$ operations
- Alternative: keep a list of “extents” which is the address of a free block and the number of consecutive free blocks starting there

Disk Scheduling Algorithms

We can minimize seek time by minimizing the distance the read/write head has to move in order to service the incoming requests.

Given a sequence of cylinders that must be visited to service a set of pending disk read/write requests, the system can order the requests to minimize seek time.

This may be done by the disk, the hardware controller, or by the operating system.

We will compare algorithms by examining their performance on a given *request queue*.

Given a disk with 200 cylinders (0–199), suppose we have 8 pending requests:

98, 183, 37, 122, 14, 124, 65, 67

and that the read/write head is currently at cylinder 53.

First-Come First-Served (FCFS)

We include this analog of FCFS CPU scheduling or FIFO page replacement mainly for comparison purposes.

Requests are serviced in queue order, for a total of 640 cylinders of movement.

Shortest Seek Time First (SSTF)/Closest Cylinder Next

Service the request next that has the shortest movement from the current position.

This is the analog of SJF CPU scheduling and OPT page replacement, but unlike those, it’s feasible here, since we do have an actual request queue (some “future knowledge”) available to us.

In our example:

65, 67, 37, 14, 98, 122, 124, 183

The total seek distance is 236 cylinders.

Potential problem: if many requests keep arriving near where the disk head is positioned, distant requests may be starved.

SCAN or Elevator Algorithm

When an elevator is going in one direction, it stops at all the floors where there is a pending request. Then it reverses direction and does the same thing.

With this algorithm, the disk arm does just this. Service requests in one direction, then reverse direction.

In our example, assuming we are “going down” at the start:

37, 14, (0), 65, 67, 98, 122, 124, 183

236 cylinders again. It is a coincidence that this is the same as SSTF.

Note that the disk arm went all the way to 0, even though there were no requests below 14. This is because this particular algorithm doesn't look ahead, it just moves back and forth from one end to the other.

We can take care of that extra movement down to 0 with ...

LOOK Algorithm

It's the same as SCAN, but the head reverses direction as soon as there are no pending requests in the current direction.

The movement is the same as SCAN, just without that move from 14 to 0 and then up to 65. This reduces the movement to 208 cylinders.

Both SCAN and LOOK can lead to non-uniform waiting times. A request near one end of the disk sometimes needs to wait for two sweeps across the disk, while other times it will be serviced very quickly. Requests near the middle have a more uniform average waiting time.

Circular Algorithms

This problem can be addressed using circular versions of SCAN (C-SCAN) and LOOK (C-LOOK), where when the disk arm gets to the end of the disk, it jumps immediately back to the other end.

Assuming the disk services requests only when “going up”, our example using these algorithms are served in order:

65, 67, 98, 122, 124, 183, 14, 37

With C-SCAN, the head goes all the way to 199 and all the way back to 0, giving total movement of 382. With C-LOOK, we do not need to go up past 183 or down past 14, making the movement total 322.

The penalty of the movement all the way back in the other direction may not be as large as it seems. Think of the mechanics of the situation – starting and stopping the disk arm takes more time than simply sweeping all the way across with just one acceleration and deceleration.

Comparing Disk Scheduling Algorithms

- SSTF or LOOK are often reasonable for a default algorithm
- SCAN and C-SCAN are better for heavily loaded systems where LOOK is unlikely to save much and SSTF runs the risk of starvation
- performance depends on the frequency and types of requests
- we may want to consider some of this when thinking about how to organize file systems

FreeBSD's ufs filesystem (the default for FreeBSD) uses an elevator algorithm. Here is the comment at the top of file `/sys/ufs/ufs/ufs_disksubr.c`:

```
/*
 * Seek sort for disks.
 *
 * The buf_queue keep two queues, sorted in ascending block order. The first
 * queue holds those requests which are positioned after the current block
 * (in the first request); the second, which starts at queue->switch_point,
 * holds requests which came in after their block number was passed. Thus
 * we implement a one way scan, retracting after reaching the end of the drive
 * to the first request on the second queue, at which time it becomes the
 * first queue.
 *
 * A one-way scan is natural because of the way UNIX read-ahead blocks are
 * allocated.
 */
```

RAID

To this point, we have talked about “partitions” as subdivisions of an individual disk. It is also possible to have a logical “partition” span multiple disks, and to create a filesystem within that logical partition.

RAID – Redundant Array of Independent/Inexpensive Disks

- multiple disks to provide reliability through redundancy
- efficiency – work can be spread across a number of disks or even disk controllers
- convenience of one large partition instead of many small ones

My experience with RAID: the former bullpen cluster:

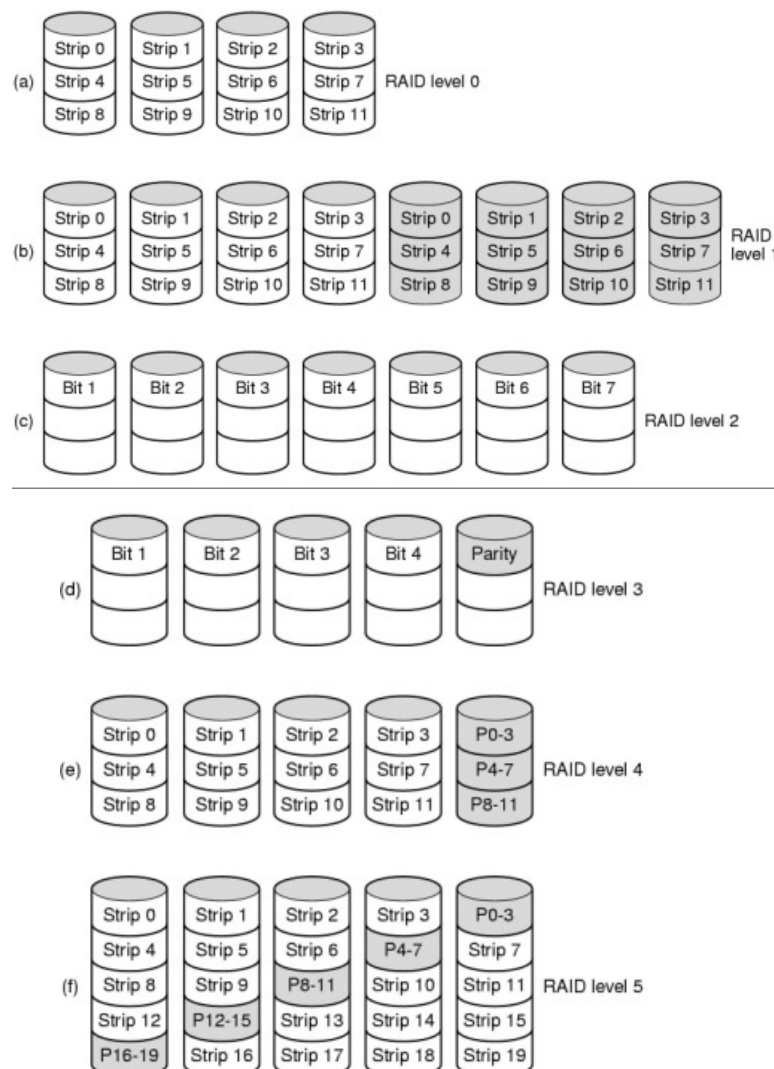
12 disks, 18 GB each. Connected to one Wide-SCSI controller. The system sees it as one big partition:

```
> df -k /export/raid
```

Filesystem	kbytes	used	avail	capacity	Mounted on
/dev/dsk/clt5d0s6	191175687	25709437	163554494	14%	/export/raid

Yes, at the time it was a big (and expensive) deal to have 191GB of space. It was 2001. Times have changed.

There are many ways to organize a RAID (Tanenbaum, Figure 5-19):



- RAID Level 0

- basically just paste together a bunch of disks to see them as one big partition
- striping
- not really a RAID, as it is not redundant
- reliability: one disk failure results in potential loss of entire partition! Actually lower the MTBF (mean time between failures)
- little or no overhead on writes
- 100% of disk space is usable for storage
- RAID Level 1
 - mirroring
 - reliability: any failed disk can be reconstructed from its mirror with a simple copy
 - all writes must be written to two disks
 - reads can come from either of two disks – spread out the load
 - 50% of disk space is usable for storage
- RAID Level 2
 - *memory-style error-correcting-code (ECC) organization*
 - use 7 synchronized disks to store 4 disks worth of data
 - for each 4 bits, compute a 7-bit Hamming-coded (see <http://mathworld.wolfram.com/HammingCode.html> word
 - Hamming codes are self-correcting for one error, can detect two errors
 - 57.1% of disk space is usable for storage
- RAID Level 3
 - *bit-interleaved parity organization*
 - like RAID-2, but use a single parity bit
 - still can recover a lost disk using the parity bit
 - two lost disks means entire partition is lost
 - can work with any number of disks
 - can include multiple parity disks
 - space overhead of the parity disk(s)
- RAID Level 4
 - *block-interleaved parity organization*
 - use stripes for parity unit, allowing disks to work independently

- still can recover a lost disk
 - parity disk can be a bottleneck as all writes require a write to the parity disk
 - space overhead of the parity disk(s)
- RAID Level 5
 - *block-interleaved distributed parity*
 - like RAID-4, but parity bit is distributed across all disks
 - This is what was used in my 2001 cluster and it actually worked when disks failed: filesystems could be accessed in “degraded” mode after disk failure, and filesystem was rebuilt automatically when the replacement disk was installed
- RAID Level 6
 - Like RAID-5 but with error correcting codes
- RAID Level 0+1, 1+0
 - Combine RAID Level 0 (for striping/efficiency) with RAID Level 1 (for redundancy)

Where does this happen?

- RAID controller – work is done by hardware, OS sees a single drive
- kernel/device driver – work is done by the OS in software – OS uses disks independently, but presents them to “users” as a single unit

In FreeBSD see “vinum” and in Linux see <https://raid.wiki.kernel.org/>

Performance Optimization

Disk Cache

Caching is an important optimization for disk accesses.

A disk cache may be located:

- main memory
- disk controller
- internal to disk drive

For a lecture assignment, you will read about a strategy used by many Unix variants to use main memory as a disk cache: the *buffer cache*.

Safety and Recovery

When a disk cache is used, there could be data in memory that has been “written” by programs, which has not yet been physically written to the disk. This can cause problems in the event of a system crash or power failure.

If the system detects this situation, typically on bootup after such a failure, a *consistency checker* is run. In Unix, this is usually the `fsck` program, and in Windows, `scandisk` or some variant. This checks for and repairs, if possible, inconsistencies in the filesystem.

Demo: `fsck`

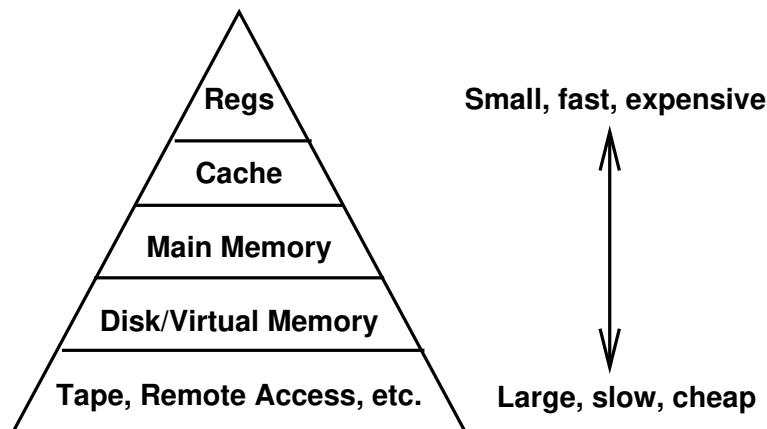
Journaling Filesystems

One way to avoid data loss when a filesystem is left in an inconsistent state is to move to a *log-structured* or *journaling* filesystem.

- record updates to the filesystem as *transactions*
- transactions are written immediately to a log and the action is committed by the OS (the application can continue), though the actual filesystem may not yet be updated
- transactions in the log are asynchronously applied to the actual filesystem, at which time the transaction is removed from the log
- if the system crashes, any pending transactions can be applied to the filesystem – main benefits are less chance of significant inconsistencies, and that those inconsistencies can be corrected from the unfinished transactions, avoiding the long consistency check
- Examples:
 - ReiserFS, see <http://en.wikipedia.org/wiki/ReiserFS>, a linux journaling filesystem
 - ext3 and ext4, most common filesystems for Linux today
 - jfs, see <http://jfs.sourceforge.net/>, IBM journaling filesystem, available for AIX, Linux
 - Related idea in FreeBSD’s filesystem: Soft Updates, see <http://www.freebsd.org/doc/en/articles/gjournal-desktop/configure-journal.html>
 - Journaling extensions to Macintosh HFS disks
 - NTFS does some journaling, but some claim it is not “fully journaled”
- the term “journaling” may also refer to systems that maintain the transaction log for a longer time, giving the ability to “undo” changes and retrieve a previous state of a filesystem

Hierarchical Storage

Recall our memory hierarchy:



Just as virtual memory uses disks to simulate a larger main memory, tapes and other removeable media can be used to simulate a larger disk.

- extend the filesystem
- small and frequently-used files remain on disk
- large, old, rarely-used files are archived on tapes
- when one of the old files is requested, the file is brought back onto the disk from the appropriate tape
- usually implemented as a jukebox of tapes or removeable disks
- tape latency is typically 1000 times that of a disk
- add in a tape robot that has to go fetch a tape and it is even worse
- or worse yet, a human who has to be notified, go to the “tape room”, find the tape, bring it to the drive, load it

These systems are found at large supercomputing centers.

- HPSS – High Performance Storage System, see <http://www.hpss-collaboration.org/>
- An older example was called UniTree

Other issues:

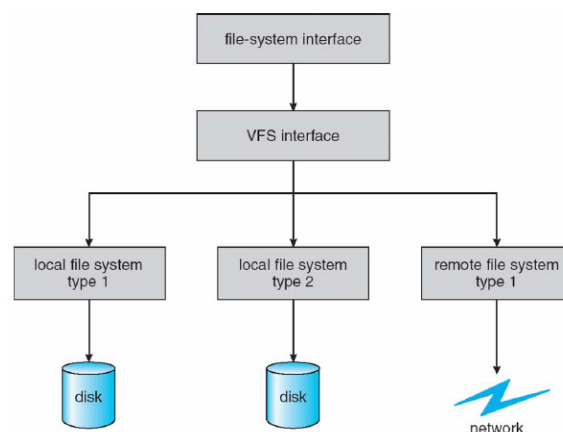
- how to decide when to archive to tape
 - retrieval from archive may be fully automated or users may need to explicitly request files from tapes
 - duplicate tapes? – tapes can be unreliable
 - when is this worthwhile? Can't we just archive information manually?
-

Virtual File System Layer

We have seen that there are many ways to organized file systems, including things like RAID and hierachical file systems. We will soon see more about networked file systems. But as users or programmers of modern computers, we rarely if ever are aware of all of these differences. We can access files exactly the same way whether the files are on a local disk with a FAT filesystem, a local disk with a ufs or ext3 filesystem, as part of a RAID, or located on a network file server. This is made possible by the *virtual file system (VFS) layer*.

VFS provide an object-oriented way of implementing and utilizing file systems.

The same system call interface (API) can be used for many different types of file systems. Even those that do not yet exist when code is written. All file access is through the API meeting the VFS interface, rather than any specific type of file system.



The VFS gets requests for various file operations from programs. It then determines, based on the path, which filesystem should be accessed. Each supported filesystem type has implementations of each operation that are aware of the details of that filesystem.

The text has a bit more information about the Linux VFS system, but our main goal is to understand the idea and why it is essential to provide convenient access to file systems.