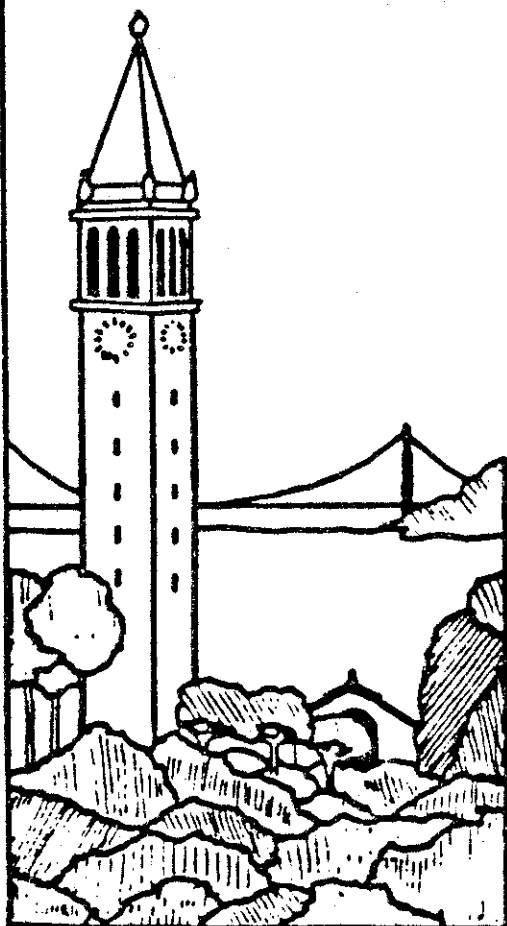


Caching in the Sprite Network File System

Michael Nelson, Brent Welch and John Ousterhout



Report No. UCB/CSD 87/345

March 1987

Computer Science Division (EECS)
University of California
Berkeley, California 94720

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE MAR 1987		2. REPORT TYPE		3. DATES COVERED 00-00-1987 to 00-00-1987	
4. TITLE AND SUBTITLE Caching in the Sprite Network File System				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) University of California at Berkeley, Department of Electrical Engineering and Computer Sciences, Berkeley, CA, 94720				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT The Sprite network operating system uses large main-memory disk block caches to achieve high performance in its file system. It provides non-write-through file caching on both client and server machines. A simple cache consistency mechanism permits files to be shared by multiple clients without danger of stale data. In order to allow the file cache to occupy as much memory as possible, the file system of each machine negotiates with the virtual memory system over physical memory usage and changes the size of the file cache dynamically. Benchmark programs indicate that client caches allow diskless Sprite workstations to perform within 5 percent of workstations with disks. In addition, client caching reduces server loading by 50% and network traffic by 75%.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 21	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Caching in the Sprite Network File System

Michael Nelson
Brent Welch
John Ousterhout

Computer Science Division
Electrical Engineering and Computer Sciences
University of California
Berkeley, CA 94720

Abstract

The Sprite network operating system uses large main-memory disk block caches to achieve high performance in its file system. It provides non-write-through file caching on both client and server machines. A simple cache consistency mechanism permits files to be shared by multiple clients without danger of stale data. In order to allow the file cache to occupy as much memory as possible, the file system of each machine negotiates with the virtual memory system over physical memory usage and changes the size of the file cache dynamically. Benchmark programs indicate that client caches allow diskless Sprite workstations to perform within 5 percent of workstations with disks. In addition, client caching reduces server loading by 50% and network traffic by 75%.

1. Introduction

Caches have been used in many operating systems to improve file system performance. Typically, caching is implemented by retaining in main memory a few of the most recently accessed disk blocks (e.g., UNIX [THOM78]). Repeated accesses to a block in the cache can be handled without involving the disk, which has two advantages. First, caching reduces delays: a block in the cache can usually be returned to a waiting process five to ten times more quickly than one that must be fetched from disk. Second, caching reduces contention for the disk arm, which may be advantageous if several processes are attempting to access files on the same disk. Measurements of timesharing systems indicate that even small caches provide substantial benefits, and that the benefits are increasing as larger physical memories permit larger caches [LEFF84, OUST85].

This paper describes a simple distributed mechanism for caching files among a networked collection of workstations. We have implemented it as part of the Sprite operating system. In Sprite, file information is cached in the main memories of both servers (workstations with disks), and clients (workstations wishing to access files on non-local disks), as shown in Figure 1. On machines with disks, the caches achieve the same effects described above, namely to reduce disk-related delays and contention. On clients, the caches also reduce the communication delays that would otherwise be required to fetch blocks from servers. In addition, client caches reduce contention for the network and for the server machines. Since server CPUs appear to be the bottleneck in several existing network file systems [SATY85, LAZO86], client caching offers the possibility of greater system scalability as well as increased performance.

There are two unusual aspects to the Sprite caching mechanism. The first is that Sprite guarantees workstations a consistent view of the data in the file system, even when multiple workstations access the same file simultaneously and the file is cached in several places at once. This is done through a simple cache consistency mechanism that flushes portions of caches and disables caching for files undergoing read-write sharing. The

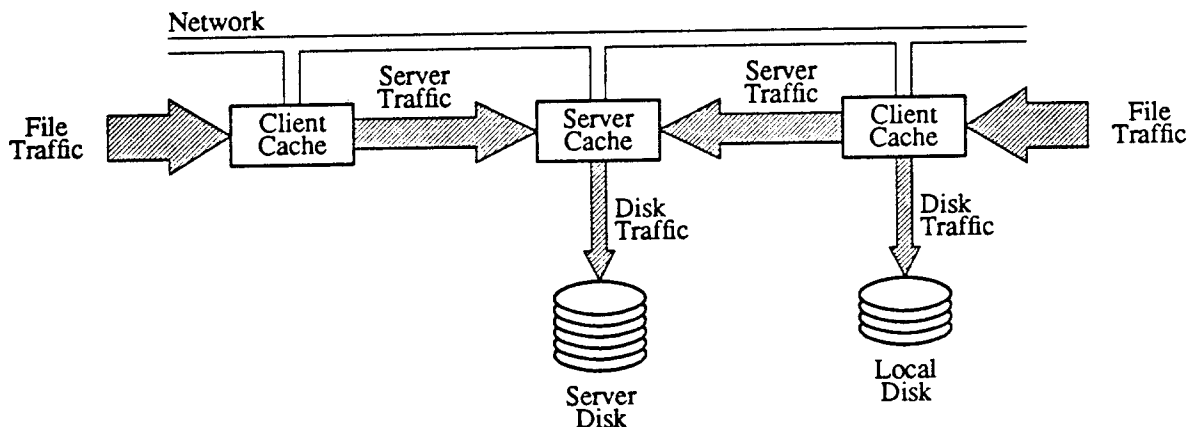


Figure 1. File caches in the Sprite system. When a process makes a file access, it is presented first to the cache of the process's workstation ("file traffic"). If not satisfied there, the request is passed either to the local disk, if any ("disk traffic"), or to the server where the file is stored ("server traffic"). Servers also maintain caches in order to reduce their disk traffic.

result is that file access under Sprite has exactly the same semantics as if all of the processes on all of the workstations were executing on a single timesharing system.

The second unusual feature of the Sprite caches is that they vary in size dynamically. This was a consequence of our desire to provide very large client caches, perhaps occupying most of the clients' memories. Unfortunately, large caches may occasionally conflict with the needs of the virtual memory system, which would like to use as much memory as possible to run user processes. Sprite provides a simple mechanism through which the virtual memory system and file system of each workstation negotiate over the machine's physical memory. As the relative needs of the two systems change, the file cache changes in size to provide the best performance to user processes.

We used a collection of benchmark programs to measure the performance of the Sprite file system, both in absolute terms and relative to Sun's Network File System [SAND85]. On average, client caching resulted in a speedup of about 10-15% for programs running on diskless workstations, relative to diskless workstations without client caches. With client caching enabled, diskless workstations completed the benchmarks only 1-5% more slowly than workstations with disks. Client caches reduced the server utilization from about 5-15% per active client to only about 1-7% per active client. Since normal users are rarely active, our measurements suggest that a single server should be able to support 50-100 clients. The benchmark programs typically executed 10-40% faster under Sprite than under NFS.

The rest of the paper is organized as follows: Section 2 gives a brief overview of Sprite; Section 3 describes prior work that motivated our cache design; Section 4 presents the basic structure of the Sprite caches; Section 5 describes the consistency protocols and Section 6 discusses the mechanism for varying the cache sizes; Section 7 presents the benchmark results; and Section 8 describes work still to be done in the areas of recovery and allocation.

2. Overview of Sprite

Sprite is a new operating system being implemented at the University of California at Berkeley as part of the development of SPUR, a high-performance multiprocessor workstation [HILL86]. A preliminary version of Sprite is currently running on Sun-2 and Sun-3 workstations, which have about 1-2 MIPS processing power and 4-16 Mbytes of main memory. The system is targeted for workstations like these and newer models likely to become available in the near future, such as SPURs; we expect the future machines to have at least five to ten times the processing power and main memory of our current machines, as well as small degrees of multiprocessing. We hope that Sprite will be suitable for networks of up to a few hundred of these workstations. Because of economic and environmental factors, most workstations will not have local disks; instead, large fast disks will be concentrated on a few server machines.

The interface that Sprite provides to user processes is much like that provided by UNIX [RITC74]. The file system appears as a single shared hierarchy accessible equally by processes on any workstation in the network (see [WELCH86a] for information on how the name space is managed). The user interface to the file system is through UNIX-like system calls such as open, close, read, and write.

Although Sprite appears similar in function to UNIX, we have completely re-implemented the kernel in order to provide better network integration. In particular, Sprite's implementation is based around a simple kernel-to-kernel remote-procedure-call (RPC) facility [WELCH86b], which allows kernels on different workstations to request services of each other using a protocol similar to the one described by Birrell and Nelson [BIRR84]. The Sprite file system uses the RPC mechanism extensively for cache management.

3. Background Work

The main motivation for the Sprite cache design came from a trace-driven analysis of file activity in several UNIX 4.2 BSD systems (hereinafter referred to as "the BSD study") [OUST85]. For those systems (three timeshared VAX-11/780s running program development, text processing, and computer-aided design applications) the BSD study showed that even small file caches are effective in reducing disk traffic, and that large caches (4-16 megabytes) work even better, cutting disk traffic by as much as 90%. For the kinds of applications measured in the BSD study it appears that increases in memory sizes will soon make it possible to keep entire file working sets in main memory, with disks serving primarily as backup devices. Although the BSD study was based on time-sharing machines rather than networks of personal workstations, we hypothesized that the results would apply in a network environment too, and that the overheads associated with remote file access could be reduced by caching on clients as well as servers (Sections 5.3 and 7 provide simulation and measurement data to support this hypothesis).

An additional motivating factor for us was a concern about server contention. A study of remote file access by Lazowska et al. concluded that the server CPU is the primary bottleneck that limits system scalability [LAZO86]. Independently, the designers of the Andrew file system decided to redesign their system in order to offload the servers [SATY85]. These experiences, plus our own informal observations of our computing environment, convinced us that client caching could substantially increase the scalability of the system.

4. Basic Cache Design

This section describes the basic organization of file caches in Sprite, and addresses three issues:

- Where should client caches be kept: main memory or local disk?
- How should caches be structured and addressed?
- What policy should be used for writing blocks back to disk?

The issues of maintaining cache consistency and varying the sizes of caches are discussed separately in the following two sections.

4.1. Caches on Disk or in Main Memory?

In several previous network filesystems (e.g. Andrew [MORR86, SATY85] and Cedar [SCHR85]), clients' file caches were kept on their local disks. For Sprite we chose to cache file data in main memory, for four reasons. First, main-memory caches permit workstations to be diskless. Second, data can be accessed more quickly from a cache in main memory than a cache on disk. Third, physical memories on client workstations are

already large enough to provide high hit ratios (e.g. a 1-Mbyte client cache provides greater than 80% read hits). As memories get larger, main-memory caches will grow to achieve even higher hit ratios. Fourth, the server caches will be in main memory regardless of where client caches are located: by using main-memory caches on clients too, we were able to build a single caching mechanism for use by both servers and clients.

4.2. Cache Structure

The Sprite caches are organized on a block basis using a fixed block size of 4 kbytes. We made this choice largely for simplicity and are prepared to revise it after we have more experience with the system. Cache blocks are addressed virtually, using a unique file identifier provided by the server and a block number within the file. We used virtual addresses instead of physical disk addresses so that clients could create new blocks in their caches without first contacting a server to find out their physical locations. Virtual addressing also allows blocks in the cache to be located without traversing the file's disk map.

For files accessed remotely, client caches hold only data blocks. Servers also cache file maps and other disk management information. These blocks are addressed in the cache using the blocks' physical disk addresses along with a special "file identifier" corresponding to the physical device.

4.3. Writing Policy

The policy used to write dirty blocks back to the server or disk has a critical effect on the system's performance and reliability. The simplest policy is to write data through to disk as soon as it is placed in any cache. The advantage of write-through is its reliability: little information is lost when a client or server crashes. However, this policy requires each write access to wait until the information is written to disk, which results in poor write performance. Since about 1/3 of all file accesses are writes [OUST85], a caching scheme based on write-through cannot reduce disk/server traffic by more than 2/3.

An alternate write policy is to delay write-backs: all blocks are written to the cache and then written through to the disk or server some time later. This policy has two advantages over write-through. First, since writes are to the cache, write accesses complete much more quickly. Second, data may be deleted before it is written back, in which case it need never be written at all. In the BSD study, 20 to 30 percent of new data was deleted within 30 seconds, and 50 percent was deleted within 5 minutes. Thus, a policy that delays writes several minutes can substantially reduce the traffic to the server or disk. Unfortunately, delayed-write schemes introduce reliability problems, since unwritten data will be lost whenever a server or client crashes.

For Sprite, we chose a delayed-write policy similar to the one used in UNIX: every 30 seconds, all dirty blocks that haven't been modified in the last 30 seconds are written back. A block written on a client will be written to the server's cache in 30-60 seconds, and will be written to disk in 30-60 more seconds. This policy avoids delays when writing files and permits modest reductions in disk/server traffic, while limiting the damage that can occur in a crash. We plan to experiment with longer write-back intervals in the future.

Another alternative is to write data back to the server when the file is closed. This approach is used in the Andrew system and NFS. Unfortunately, the BSD study found that 75 percent of files are open less than 0.5 seconds and 90 percent are open less than 10 seconds. This indicates that a write-on-close policy will not significantly reduce disk or server traffic. In addition, the write-on-close policy requires the closing process to delay while the file is written through, which reduces the performance advantages of delayed writes.

5. Cache Consistency

Allowing clients to cache files introduces a consistency problem: what happens if a client modifies a file that is also cached by other clients? Can subsequent references to the file by other clients return "stale" data? Most existing network file systems provide only limited guarantees about consistency. For example, the NFS and Andrew systems guarantee that once a file is closed all data is back on the server and future opens by other clients will cause their caches to be updated to contain the new version. Under conditions of "sequential write-sharing" (a file is shared but is never open simultaneously for reading and writing on different clients), each client will always see the most up-to-date version of the file. However, if a file in NFS or Andrew is open simultaneously on several clients and one of them modifies it, the other clients will not see the changes immediately; users are warned not to attempt this kind of sharing (which we call "concurrent write-sharing").

For Sprite we decided to permit both concurrent and sequential write-sharing. Sprite guarantees that whenever a process reads data from a file, it receives the most recently written data, regardless of when and where the data was last written. We did this in order to make the user view of the file system as clean and simple as possible, and to encourage use of the file system as a shared system-wide store for exchanging information between different processes on different machines. We hope that shared files will be used to simplify the implementation of system services such as print spoolers and mailers. Of course, we still expect that concurrent write-sharing will be infrequent, so the consistency algorithm is optimized for the case where there is no sharing.

The only other network file system we know of that permits concurrent write-sharing is Locus [POPEK85]. It uses a complex mechanism based on passing tokens between the workstations that are accessing the file. For Sprite, we adopted a simpler approach that uses the servers as centralized control points for cache consistency. Each server guarantees cache consistency for all the files on its disks, and clients deal only with the server for a file: there are no direct client-client interactions. The following subsections deal separately with the problems of concurrent write-sharing and sequential write-sharing.

The Sprite algorithm depends on the fact that the server is notified whenever one of its files is opened or closed, so it can detect when concurrent write-sharing is about to occur. This approach prohibits performance optimizations (such as name caching) that cause clients to open files without contacting the files' servers. The benchmark results of Section 7 suggest that such optimizations would only provide small additional performance improvements.

5.1. Concurrent Write-Sharing

Concurrent write-sharing occurs for a file when it is open on multiple clients and at least one of them has it open for writing. Sprite deals with this situation by disabling client caching for the file, so that all reads and writes for the file go through to the server. When a server detects (during an "open" operation) that concurrent write-sharing is about to occur for a file, it takes two actions. First, it notifies the client that has the file open for writing, if any, telling it to write all dirty blocks back to the server. There can be at most one such client. Second, the server notifies all clients that have the file open, telling them that the file is no longer cacheable. This causes the clients to remove all of the file's blocks from their caches. Once these two actions are taken, clients will send all future accesses for that file to the server. The server's kernel serializes the accesses to its cache, producing a result identical to running all the client processes on a single timeshared machine.

Caching is disabled on a file-by-file basis, and only when concurrent write-sharing occurs. A file can be cached simultaneously by many clients as long as none of them is writing the file, and a writing client can cache the file as long as there are no concurrent readers or writers on other workstations. When a file becomes non-cacheable, only those clients with the file open are notified; if other clients have some of the file's data in their caches, they will take consistency actions the next time they open the file, as described below. A non-cacheable file does not become cacheable again until it has been closed on all clients.

5.2. Sequential Write-Sharing

Sequential write-sharing occurs when a file is modified by one client, closed, then opened by some other client. There are two potential problems associated with sequential write-sharing. First, when a client opens a file it may have out-of-date blocks in its cache. To solve this problem, servers keep a version number for each file, which is incremented each time the file is opened for writing. Each client keeps the version numbers of all the files in its cache. When a file is opened, the client compares the server's version number for the file with its own. If they differ, the client flushes the file from its cache.

The second potential problem with sequential write-sharing is that the current data for the file may be in some other client's cache (the last writer need not have flushed dirty blocks back to the server when it closed the file). Servers handle this situation by keeping track of the last writer for each file; this client is the only one that could potentially have dirty blocks in its cache. When a client opens a file the server notifies the last writer (if there is one and if it is a different client than the opening client), and waits for it to write its dirty blocks through to the server. This ensures that the reading client will receive up-to-date information when it requests blocks from the server.

5.3. Simulation Results

We used the trace data from the BSD study to estimate the overheads associated with cache consistency, and also to estimate the overall effectiveness of client caches. The data were used as input to a simulator that treated each timesharing user as a separate client workstation in a network with a single file server. The results are shown in Table 1. Client caching reduced server traffic by over 70% and resulted in read hit ratios of more than 80%.

Server Traffic With Cache Consistency				
Client Cache Size	Blocks Read	Blocks Written	Total	Traffic Ratio
0 Mbyte	445815	172546	618361	100%
0.5 Mbyte	102469	96866	199335	32%
1 Mbyte	84017	96796	180813	29%
2 Mbytes	77445	96796	174241	28%
4 Mbytes	75322	96796	172118	28%
8 Mbytes	75088	96796	171884	28%

Table 1. Client caching simulation results, based on trace data from BSD study. Each user was treated as a different client, with client caching and a 30-second delayed-write policy. The table shows the number of read and write requests made by client caches to the server, for different client cache sizes. The "Traffic Ratio" column gives the total server traffic as a percentage of the total file traffic presented to the client caches. Write-sharing is infrequent: of the write traffic, 4041 blocks were written through because of concurrent write-sharing and 6887 blocks were flushed back because of sequential write-sharing.

Server Traffic, Ignoring Cache Consistency				
Client Cache Size	Blocks Read	Blocks Written	Total	Traffic Ratio
0 Mbyte	445815	172546	618361	100%
0.5 Mbyte	80754	93663	174417	28%
1 Mbyte	52377	93258	145635	24%
2 Mbytes	41767	93258	135025	22%
4 Mbytes	38165	93258	131423	21%
8 Mbytes	37007	93258	130265	21%

Table 2. Traffic without cache consistency. Similar to Table 1 except that cache consistency issues were ignored completely.

Table 2 presents similar data for a simulation where no attempt was made to guarantee cache consistency. A comparison of Tables 1 and 2 shows that about 25% of all server traffic is due to cache consistency, and that most of the cache consistency overhead is from blocks that are flushed from client caches and must be re-loaded. Table 2 is not realistic, in the sense that it simulates a situation where incorrect results would have been produced; nonetheless, it provides an upper bound on the cache consistency overheads.

6. Virtual Memory and the File System

In addition to guaranteeing coherency between the client caches, we wanted to permit each client cache to be as large as possible. For example, applications that don't require much virtual memory should be able to use most of the physical memory as a file cache. However, if the caches were fixed in size (as they are in UNIX), then large caches would leave little physical memory for running user programs, and it would be difficult to run applications with large virtual memory needs. In order to get the best overall

performance, Sprite allows each file cache to grow and shrink dynamically in response to changing demands on the machine's virtual memory system and file system. This is accomplished by having the two modules negotiate over physical memory usage.

The file system module (FS) and the virtual memory module (VM) each manage a separate pool of physical memory pages. VM keeps its pages in approximate LRU order through a version of the clock algorithm [NELS86]. FS keeps its cache blocks in perfect LRU order since all block accesses are through the "read" and "write" system calls. Each system keeps a time-of-last-access for each page or block. Whenever either module needs additional memory (because of a page fault or a miss in the file cache), it compares the age of its oldest page with the age of the oldest page from the other module. If the other module has the oldest page, then it is forced to give up that page; otherwise the module recycles its own oldest page.

The approach just described has two potential problems: double-caching and multi-block pages. Double-caching can occur because VM is a user of the file system: backing storage is implemented using ordinary files, and read-only code is demand-loaded directly from executable files. A naive implementation might cause pages being read from backing files to end up in both the file cache and the VM page pool; pages being eliminated from the VM page pool might simply get moved to the file cache, where they would have to age for another 30 seconds before being sent to the server. To avoid these inefficiencies, the VM system bypasses the local file cache when reading and writing backing files. A similar problem occurs when demand-loading code from its executable file. In this case, the pages may already be in the file cache (e.g. because the program was just recompiled). If so, the page is copied to the virtual memory page pool and the block in the file cache is given an "infinite" age so that it will be replaced before anything else in memory.

Although VM bypasses its local file cache when reading and writing backing files, the backing files *will* be cached on servers. This makes servers' memories into an extended main memory for their clients.

The second problem with the negotiation between VM and FS occurs when virtual memory pages are large enough to hold several file blocks. Is the LRU time of a page in the file cache the age of the oldest block in the page, the age of the youngest block in the page, or the average age of the blocks in the page? Once it is determined which page to give back to virtual memory, what should be done with the other blocks in the page if they have been recently accessed? For our Sun-3 implementation of Sprite, which has 8-kbyte pages but 4-kbyte file blocks, we used a simple solution: the age of a page is the age of the youngest block in the page, and when a page is relinquished all blocks in the page are removed. We are currently investigating the effect of this policy on file system performance.

We also considered more centralized approaches to trading off physical memory between the VM page pool and the file cache. One possibility would be to access all information through the virtual memory system. To access a file, it would first be mapped into a process's virtual address space and then read and written just like virtual memory, as in Apollo's DOMAIN system [LEACH83]. This approach would eliminate the file cache entirely; the standard page replacement mechanisms would automatically balance physical memory usage between file and program information. We rejected this approach for several reasons, the most important of which is that it would have forced us

to use a more complicated cache consistency scheme. A mapped-file approach requires a file's pages to be cached in a workstation's memory before they can be accessed, so we would not have been able to implement cache consistency by refusing to cache shared files.

Another possible approach would have been to implement a centralized physical memory manager, from which both VM and FS would make page requests. The centralized manager would compute page ages and make all replacement decisions. We rejected this approach because the most logical way to compute page ages is different in VM than in FS. The only thing the two modules have in common is the notion of page age and LRU replacement. These shared notions are retained in our distributed mechanism, while leaving each module free to age its own pages in the most convenient way. Our approach also permits us to adjust the relative aging rates for VM and FS pages, if that should turn out to be desirable.

7. Benchmarks

This section describes a series of benchmarks we ran to measure the performance of the Sprite file system. Our goal was to measure the benefits provided by client caches in reducing delays and contention:

- How much more quickly can file-intensive programs execute with client caches than without?
- How much do client caches reduce the load placed on server CPUs?
- How much do client caches reduce the network load?

In addition to answering these questions for our current machines, we have tried to predict how the benefits of client caches are likely to change in the future, as CPU speeds and memory sizes increase. See Table 3 for information about the machines on which the benchmarks were run.

7.1. Micro-benchmarks

To determine the raw read and write performance of the Sprite file system, we wrote simple programs that read or write large volumes of data. Before running the programs, we rigged the system so that all the accesses would be satisfied in a particular place (e.g. the client's cache). Table 4 shows the I/O speeds achieved to and from caches and disks in different locations.

Type	CPU Speed	Memory	Disk
Sun-2/50	0.7 MIPS	4 Mbytes	None
Sun-2/120	0.7 MIPS	4 Mbytes	70 Mbyte Micropolis
Sun-3/75	2 MIPS	8 Mbytes	None
Sun-3/180	2 MIPS	16 Mbytes	400 Mbyte Fujitsu Eagle

Table 3. Machine configurations. Each benchmark was run in either a Sun-2 configuration (Sun-2/50 and Sun-2/120 clients, Sun-2/120 server) or a Sun-3 configuration (Sun-3/75 clients, Sun-3/180 server).

Read & Write Throughput, kbytes/second					
System		Local Cache	Server Cache	Local Disk	Server Disk
Sun-2	Read	1048	165	135	80
	Write	725	138	100	75
Sun-3	Read	3300	415	224	216
	Write	1905	350	205	190

Table 4. Maximum rates at which programs can read and write file data in various places, using large files accessed sequentially.

Table 4 contains two important results. First, a client can access bytes in its own cache 5-8 times faster than those in the server's cache. This means that, in the best case, client caching could permit an application program to run as much as 5-8 times faster than it could without client caching. The second important result is that a client can read and write the server's cache almost as fast as a local disk (in our measurements the server cache is faster than local disk, but Sprite's disk-reading code is not very efficient; UNIX can access a local disk 10-20% faster than Sprite can access a server cache). In the future, as CPUs get much faster but disks don't, the server's cache should become much faster than a local disk. Even for paging traffic, this suggests that a large server cache may provide better performance than a local disk.

7.2. Macro-benchmarks

The micro-benchmarks give an upper limit on the possible benefits of client caching. To see how much of this potential speedup can be achieved in real applications, we ported several widely-used programs from UNIX to Sprite and measured them under varying conditions. Table 5 describes the benchmark programs. We were surprised at

Program	Description	I/O (kbytes/sec)	
		Read	Write
Fs-make	Use the "make" program to recompile the Sprite file system: 31 source files, 28,000 lines of C source code.	15	3.8
Csh-make	Recompile csh program: 26 source files, 16,000 lines of C source code.	12	3.6
Simulator	Simulate set-associative cache memory using 1057-kbyte address trace.	9.6	0
Sort	Sort a 635-kbyte file.	9.5	9.5
Diff	Compare 2 identical 600-kbyte files.	32	0
Nroff	Format the text of this paper.	6.0	5.6

Table 5. Macro-benchmarks. The I/O columns give the average rates at which file data were read and written by the benchmark when run on Sun-2's with local disks; they measure the benchmark's I/O intensity.

how little I/O is performed even by these file-intensive programs: the most intensive of them only required about 32 kbytes/second of I/O on a Sun-2.

7.2.1. Application Speedups

Table 6 gives the total elapsed time to execute each of the macro-benchmarks with local or remote disks and with client caches enabled or disabled. Even without client caching, diskless machines were generally only about 10-20% slower than those with disks. With client caching enabled and a warm start (caches already loaded by previous activity), the difference between diskless machines and those with disks was almost unmeasurable; in the worst case, it was only about 5%. Figure 2 shows how the performance varied with the size of the client cache.

We expect the advantages of client caching to improve over time, for two reasons. First, increasing memory sizes will make larger and larger caches feasible, which will increase their effectiveness. Second, processor speeds are increasing faster than network or disk speeds; without caches, workstations will end up spending more and more of their time waiting for the network or disk.

Benchmark	Local Disk, with Cache		Diskless, Server Cache Only		Diskless, Client & Server Caches	
	Cold	Warm	Cold	Warm	Cold	Warm
Fs-make	24:15 100%	24:16 100%	30:37 126%	30:20 125%	25:47 106%	25:26 105%
Csh-make	9:01 102%	8:51 100%	10:23 118%	10:18 116%	9:15 105%	9:01 102%
Simulator	1:59 106%	1:52 100%	2:04 111%	1:57 105%	2:04 111%	1:52 100%
Sort	2:13 104%	2:08 100%	2:28 116%	2:22 111%	2:18 108%	2:09 101%
Diff	0:35 184%	0:18 100%	0:41 216%	0:24 121%	0:41 216%	0:18 100%
Nroff	2:13 102%	2:11 100%	2:25 111%	2:24 110%	2:14 102%	2:11 100%

Table 6. Execution times with and without local disks and caching, measured on Sun-2's. The top number for each run is total elapsed time, in minutes and seconds. The bottom number is normalized relative to the warm-start time with a local disk. "Cold" means that all caches, both on server and client, were empty at the beginning of the run. "Warm" means that the program was run once to load the caches, then timed on a second run. In the "Diskless, Server Cache Only" case, the client cache was disabled but the server cache was still enabled. In all other cases, caches were enabled on all machines. All caches were allowed to vary in size using the VM-FS negotiation scheme.

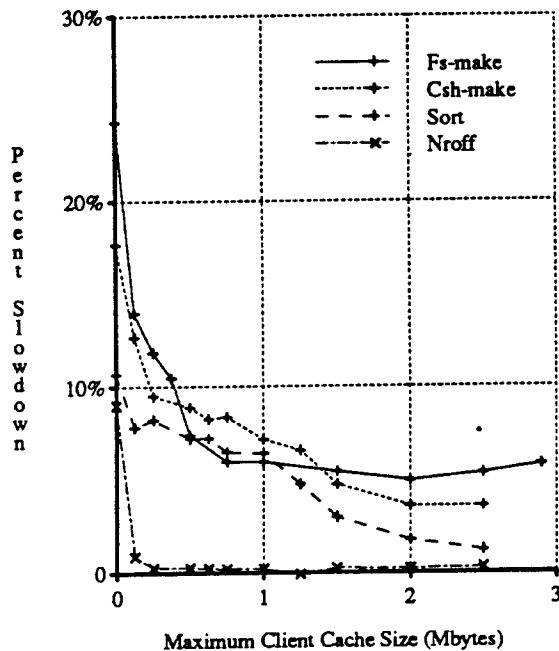


Figure 2. Client degradation (diskless Sun-2's with client caches, warm start) as a function of maximum client cache size. "Degradation" is relative to the time required to execute the benchmark with a local disk and warm cache. For each point, the maximum size of the client cache was limited to a particular value. We believe that the slight upturn when the cache size becomes unlimited is due to the overhead of negotiation between the file system and virtual memory system.

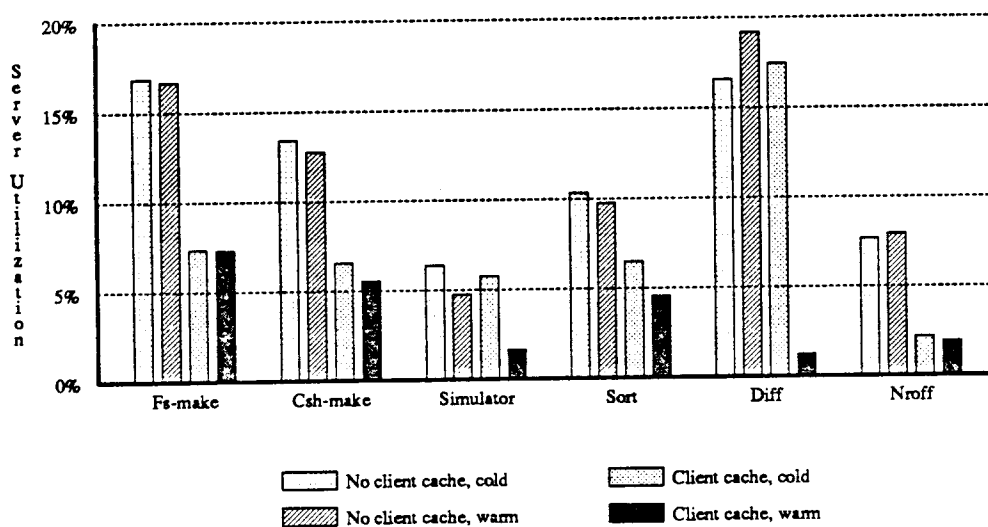


Figure 3. Client caching reduces server loading by a factor of 2-5 (measured on Sun-2's with variable-size client caches).

7.2.2. Server Contention

One of the most beneficial effects of client caching is its reduction in the load placed on server CPUs. Figure 3 shows the server CPU utilization with and without client caching. In general, a diskless client without client cache utilized about 5-20% of the server's CPU. With client caching, the server utilization dropped by a factor of two or more, to 2-7%.

We also tested the effects of contention for servers by running several versions of the Csh-make benchmark simultaneously on different clients. Each client used a different copy of the input and output files, so there was no cache consistency overhead. Figure 4 shows the effects of loading on the client speed and on the server's CPU, with and without client caches. Without client caches, there was significant performance degradation when more than a few clients were active at once. With client caches and six active clients, each ran at a speed within 10% of what it could have achieved with a local disk; server utilization in this situation was only about 30%.

The measurements of Figure 4 suggest that with client caches a single Sun-2 server can support the needs of at least ten Sun-2 clients simultaneously running file-intensive programs. However, typical users spend only a small fraction of their time running such programs. For example, the BSD study measured average I/O rates per user of .2-2 kbytes/second whereas the application shown in Figure 4 had an average I/O rate of 15 kbytes/second. These numbers suggest that one Sun-2 Sprite file server should be able to support at least 50-100 Sun-2 users, if the users are similar to those measured in the BSD

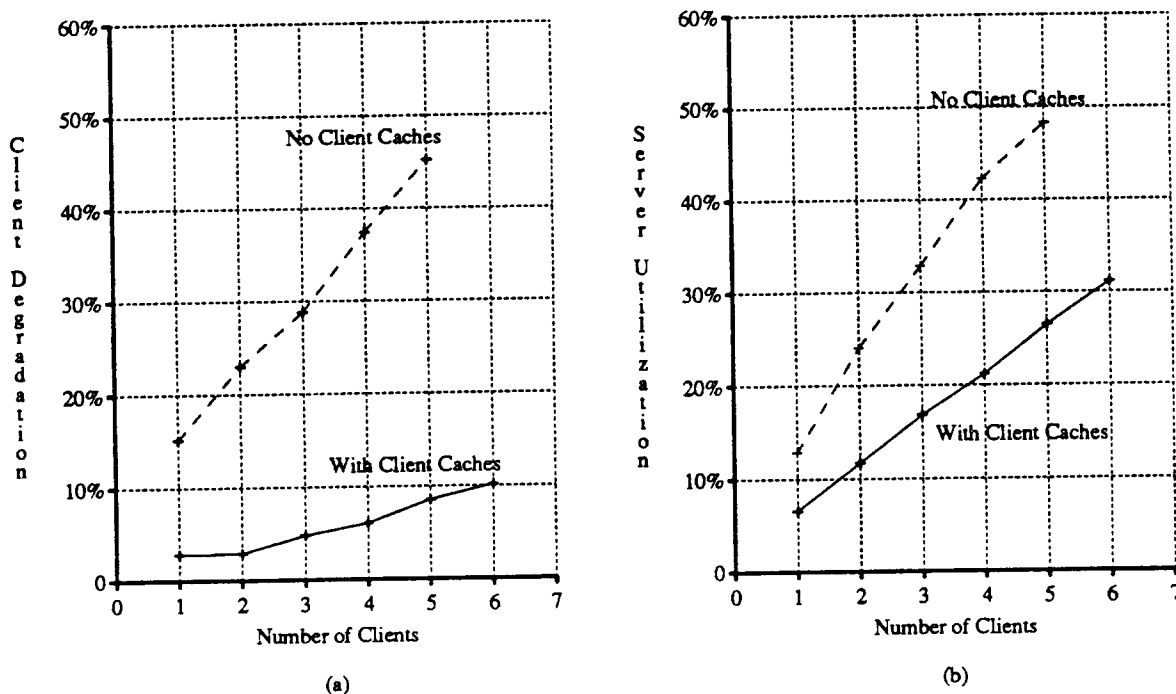


Figure 4. Effect of multiple diskless clients running the Csh-make benchmark simultaneously on different files in a Sun-2 configuration with variable-size client caches. (a) shows additional time required by each diskless client to complete the benchmark, relative to a single client running with local disk. (b) shows server utilization.

study. The server capacity should not change much with increasing CPU speeds, as long as both client and server CPU speeds increase at about the same rate. In a system with servers that are more powerful than clients, the server capacity should be even higher than this.

7.2.3. Network Contention

In their analysis of diskless file access, Lazowska et al. concluded that network loading is not a major factor in today's network file systems [LAZO86]. For today's machines, our measurements support their conclusion: even without client caching, the most intensive Sun-2 benchmark only required only about 80 kbits/second of the 10000-kbits/second of available bandwidth.

In the future, though, we expect network traffic to become more and more of an issue. The same benchmarks running on diskless Sun-3's instead of Sun-2's produced a network load of about 250 kbits/second, or 2.5% of the total Ethernet bandwidth. It seems likely that machines at least five times faster than Sun-3's will be available within a few years (e.g., the SPUR workstations under development at Berkeley); a single one of these machines would utilize 10-15% of the Ethernet bandwidth running the benchmarks without client caching. For these machines, and the even faster ones to follow, one of the main advantages of client caching is that it reduces network loading by a factor of 4 or more (see Figure 5). Without client caches, application performance may be limited by network transmission delays, and the number of workstations on a single Ethernet may be limited by the bandwidth available on the network.

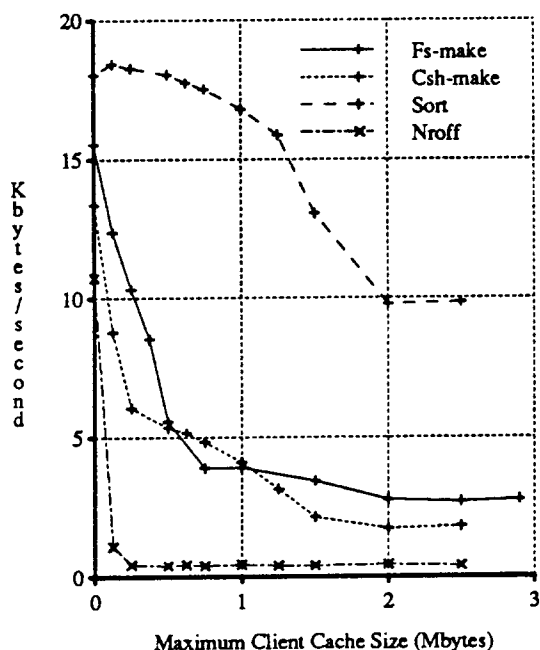


Figure 5. Network traffic (Kbytes per second, averaged over the life of the benchmark) as a function of client cache size (diskless Sun-2's, warm start). Only bits of file data were counted; bits transmitted in packet headers and control packets were not counted.

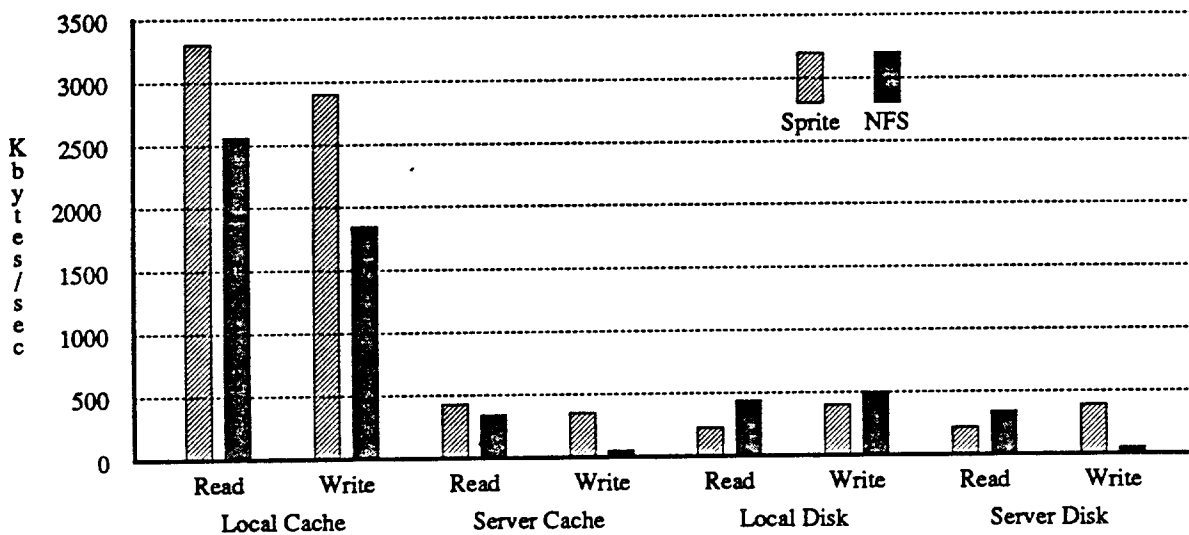


Figure 6. Raw read and write throughput for Sprite and NFS, measured on Sun-3's in the same way as Table 4. There is no way to write to the server's cache in NFS without also writing to disk, so those two numbers are the same.

Sprite vs. NFS						
Benchmark	NFS		Sprite Limited		Sprite Max	
	Local Disk	Diskless	Local Disk	Diskless	Local Disk	Diskless
Fs-make	10:51	15:52	9:12	9:54	8:54	9:31
Csh-make	3:52	4:58	3:25	3:35	3:17	3:28
Simulator	:40	:42	:40	:42	:40	:40
Sort	:44	:50	:43	:47	:41	:41
Diff	:04	:06	:05	:08	:05	:05
Nroff	:52	:53	:51	:51	:51	:51

Table 8. A comparison between Sprite and NFS. All benchmarks were run on Sun-3 configurations. The NFS numbers were measured with Sun's UNIX release 3.0, with 4k blocks. All numbers are elapsed times for warm starts, in minutes and seconds. Two Sprite cases are shown, one where the Sprite caches were restricted to be no larger than the NFS caches ("Sprite Limited"), and one where Sprite caches were allowed to grow as large as possible ("Sprite Max").

7.2.4. Sprite vs. NFS

As a final measurement of the performance of Sprite's file system, we compared it to Sun's NFS, which has become a commercial standard. See Figures 6 and 7 and Table 8. In almost all cases Sprite performs as well as NFS or better, and Sprite provides as much as 40% better performance for some benchmarks. The two systems are sufficiently different that it is hard to attribute the performances differences to any one thing, but two features of the systems appear to have a major impact on the comparison: write-through and read-ahead.

Sprite does not perform either write-through or read-ahead, while NFS does both. Write-through is used in NFS to simplify recovery after server crashes, but it results in poor write performance. The raw write bandwidth of NFS is only one-fifth of the write bandwidth in Sprite (see Figure 5), and this appears to account for most of the performance differential in the Fs-make and Csh-make applications.

The second factor is read-ahead, which NFS uses to its advantage. For example, Figure 6 shows that NFS can retrieve data from the server's disk just as quickly as it can from the server's cache; this is due to read-ahead on the server. Sprite does not currently do any read-ahead, which causes a noticeable performance degradation for applications that read large files sequentially, such as the Diff benchmark. Only when the entire file fits in the cache does Sprite perform as well on this benchmark as NFS. Although high cache hit ratios reduce the benefit of read-ahead, there still appear to be applications where read-ahead is beneficial; we plan to implement read-ahead in Sprite.

8. Future Work

There are two issues concerning client caching that we have not yet resolved in the Sprite implementation: crash recovery and disk overflow. The current system is fragile due to the amount of state kept in the main memory of each server. If a server crashes, then all the information in its memory is lost, including dirty blocks in its cache and information about open files. As a result, all client processes using files from the server usually have to be killed. In contrast, the servers in Sun's NFS are stateless. This results in less efficient operation (since all important information must be written through to disk), but it means that clients can recover from server crashes: the processes are put to sleep until the server reboots, then they continue with no ill effects.

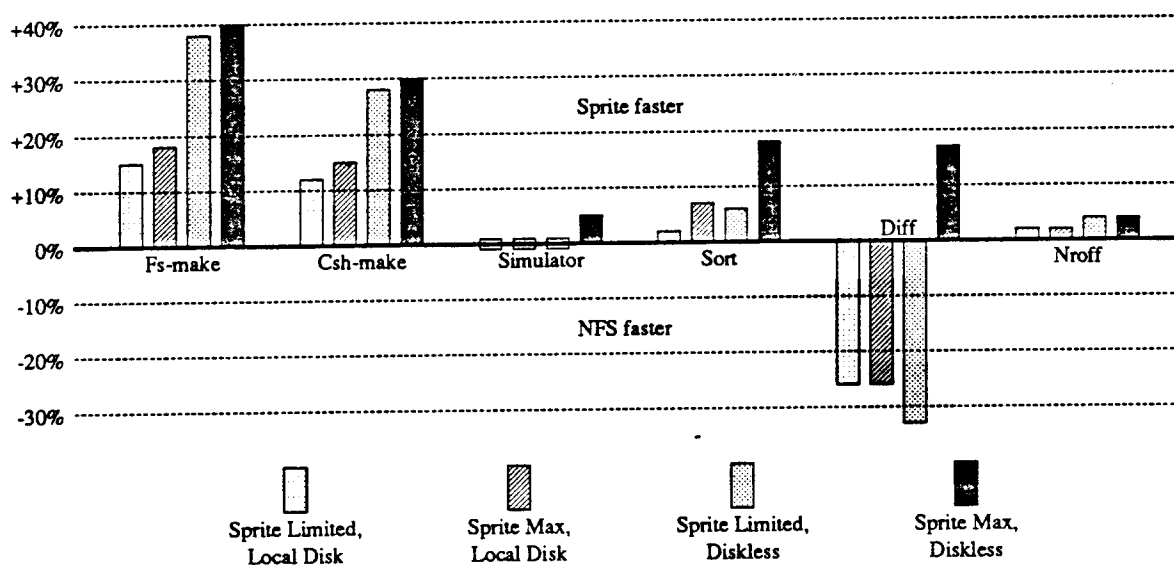


Figure 7. Sprite performance relative to NFS on Sun-3 configurations. This is the same data as in Table 8, except that Sprite's completion times are characterized as percentages better or worse than the corresponding NFS time (e.g. Sprite diskless times are compared to NFS diskless times).

We are currently exploring ways to provide better crash recovery in Sprite. It appears from our performance measurements that server caches could be made write-through without significant performance degradation. This would guarantee that no file data would be lost on server crashes. Client caches would still use a delayed-write policy, so the extra overhead of writing through the server cache would only be incurred by the background processes that clean client caches. In addition, clients should be able to provide servers with enough information to re-open files after a server crash. We hope that these two features will enable clients to continue transparently after server crashes.

The second unresolved issue has to do with "disk-full" conditions. In the current implementation, a client does not notify the server when it allocates new blocks for files. This means that when the client eventually writes the new block to the server (as much as 30 seconds later), there may be no disk space available for the block. In UNIX, a process is notified at the time of the "write" system call if the disk is full. We plan to provide similar behavior in Sprite with a simple quota system in which each client is given a number of blocks from which it can allocate disk space. If the client uses up its quota, it requests more blocks from the server. When the amount of free disk space is too small to give quotas to clients, clients will have to submit explicit disk allocation requests to the server whenever they create new blocks.

9. Conclusions

Sprite's file system demonstrates the viability of large caches for providing high-performance access to shared file data. Large caches on clients allow diskless client workstations to attain performance comparable to workstations with disks. This performance is attained while utilizing only a small portion of servers' CPU cycles. The caches can be kept consistent using a simple algorithm because write-sharing is rare. By varying the cache sizes dynamically, Sprite permits the file caches to become as large as possible without impacting virtual memory performance.

The high performance attainable with client caches casts doubts on the need for local disks on client workstations. For users considering the purchase of a local disk, our advice is to spend the same amount of money on additional memory instead. We believe that this would improve the performance of the workstation more than the addition of a local disk: it would not only improve file system performance by allowing a larger cache, but it would also improve virtual memory performance.

10. Acknowledgments

This work was supported in part by the Defense Advanced Research Projects Agency under contract N00039-85-C-0269, in part by the National Science Foundation under grant ECS-8351961, and in part by General Motors Corporation. Andrew Cheren-son, Fred Douglass, Garth Gibson, Mark Hill, Randy Katz, and Dave Patterson provided numerous helpful comments that improved the presentation of the paper.

11. References

[BIRR84]

Birrell, A.D., Nelson, B.J. "Implementing Remote Procedure Calls." *ACM Transactions on Computer Systems*, Vol. 2, No. 1, February 1984, pp. 39-59.

[HILL86]

Hill, M.D., et al. "Design Decisions in SPUR." *IEEE Computer*, Vol. 19, No. 11, November 1986, pp. 8-22.

[KLEI86]

Kleiman, S.R. "Vnodes: An Architecture for Multiple File System Types in Sun UNIX." *Proceedings of the USENIX 1986 Summer Conference*, June 1986, pp. 238-247.

[LAZO86]

Lazowska, E.D., Zahorjan, J., Cheriton, D., and Zwaenepoel, W. "File Access Performance of Diskless Workstations." *ACM Transactions on Computer Systems*, Vol. 4, No. 3, August 1986, pp. 238-268.

[LEACH83]

Leach, P.J., et al. "The Architecture of an Integrated Local Network." *IEEE Journal on Selected Areas in Communications*, Vol. SAC-1, No. 5, November 1983, pp. 842-857.

[LEFF84]

Leffler, S., Karels, M., and McKusick, M.K. "Measuring and Improving the Performance of 4.2BSD", *Proceedings of the USENIX 1984 Summer Conference*, June 1984, pp. 237-252.

[MORR86]

Morris, J.H., et al. "Andrew: A Distributed Personal Computing Environment." *Communications of the ACM*, Vol. 29, No. 3, March 1986, pp 184-201.

[NELS86]

Nelson, M. "Virtual Memory for the Sprite Operating System." Technical Report UCB/CSD 86/301, Computer Science Division (EECS), University of California, Berkeley, 1986.

[OUST85]

Ousterhout, J.K. et al. "A Trace-Driven Analysis of the 4.2 BSD UNIX File System." *Proceedings of the 10th Symposium on Operating Systems Principles*, December 1985, pp. 15-24.

[POPEK85]

G.J., Walker, B.J, editors. "The LOCUS Distributed System Architecture." The MIT Press, Cambridge, Mass., 1985.

[RITC74]

Ritchie, D.M., and Thompson, K. "The UNIX Time-Sharing System." *Communications of the ACM*, Vol. 17, No. 7, July 1974, pp. 365-375.

[SAND85]

Sandberg, R. et al. "Design and Implementation of the Sun Network Filesystem." *Proceedings of the USENIX 1985 Summer Conference*, June 1985, pp. 119-130.

[SATY85]

Satyanarayanan, M. et al. "The ITC Distributed File System: Principles and Design." *Proceedings of the 10th Symposium on Operating Systems Principles*, December 1985, pp. 35-50.

[SCHR85]

Schroeder, M.D., Gifford, D.K. and Needham, R.M. "A Caching File System for a Programmer's Workstation." *Proceedings of the 10th Symposium on Operating Systems Principles*, December 1985, pp. 25-34.

[THOM78]

Thompson, K. "UNIX Time-Sharing System: UNIX Implementation." *Bell System Technical Journal*, Vol. 57, No. 6, July-August 1978, pp. 1931-1946.

[WELCH86a]

Welch, B., Ousterhout, J. "Prefix Tables: A Simple Mechanism for Locating Files in a Distributed Filesystem." *Proceedings of the 6th International Conference on Distributed Computing Systems*, May 1986, pp. 184-189.

[WELCH86b]

Welch, B. "The Sprite Remote Procedure Call System." Technical Report UCB/CSD 86/302, Computer Science Division (EECS), University of California, Berkeley, 1986.