

AD-A135 201

THE DISTRIBUTED V KERNEL AND ITS PERFORMANCE FOR
DISKLESS WORKSTATIONS; STANFORD UNIV CA DEPT OF
COMPUTER SCIENCE D R CHERTON ET AL JUL 83
STAN CS 83 973 MDA903 80 C 0102

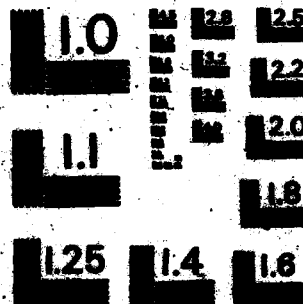
1/1

UNCLASSIFIED

1/6 17/2 NI



END
1/6



MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS-1963-A

July 1983

Report No. STAN-CS-83-973
Also numbered CSL 246

②

The Distributed V Kernel and Its Performance for Diskless Workstations

by

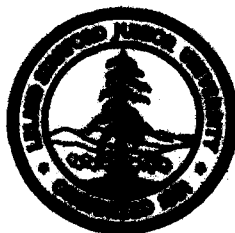
David R. Cheriton and Willy Zwanevool

Contract NDA-903-80-0-0103

Department of Computer Science

Stanford University
Stanford, CA 94305

APPROVED FOR PUBLIC RELEASE
DISTRIBUTION UNLIMITED



DTIC FILE COPY

DTIC
ELECTE
DEC 1 1983
S A D

88 : 10 04 015

AD-A185-201

The Distributed V Kernel and Its Performance for Diskless Workstations

David R. Cheriton and Willy Swenepoel

Computer Systems Laboratory
Departments of Computer Science and Electrical Engineering
Stanford University

The author

Abstract

The distributed V kernel is a microkernel-based system that provides efficient local and network transparent communication. It is primarily being used to do experiments of diskless workstations connected by a high-speed local network to a set of file servers. We describe a performance evaluation of the kernel, with particular emphasis on the cost of network file access. Our results show that over a local network:

1. Diskless workstations can access remote files with minimal performance penalty.
2. The V storage facility can be used to access remote files at comparable cost to a local disk-based workstation file access system.

It is concluded

that it is feasible to build a distributed system with all network communication using the V storage facility even when most of the network nodes have no permanent storage.

1. Introduction

The distributed V kernel is a microkernel-based system that provides efficient local and network transparent communication. The kernel consists of a system call interface, a set of system calls, and a set of system services. The system call interface is implemented in the kernel, and the system calls are implemented in the system services. The system services are implemented in the kernel, and the system call interface is implemented in the kernel.

Network interprocess communication is predominantly used for remote file access since most SUN workstations at Stanford are configured without a local disk.

This paper reports our experience with the implementation and use of the V kernel. Of particular interest are the experimental results of our approach, namely:

1. The use of diskless workstations with all secondary storage provided by behind the server.
2. The use of a general purpose network transparent communication facility for support of special-purpose file access protocols and, in particular, the use of a Thrash-like transparent communication mechanism.

The most experimental approach is to configure workstations with a small local disk, using server-based file systems for primary storage. Diskless workstations, however, have a number of advantages including:

1. Lower hardware cost per workstation.
2. Greater responsiveness and simplicity of code with shared file servers.
3. Little or no memory or processing overhead on the workstations for the kernel and disk handling.
4. Tight coupling with application, reducing the distribution of files.

The major disadvantage is the overhead of processing all file access requests. Our approach involves using the V storage facility to access remote files. The V storage facility is a microkernel-based system that provides efficient local and network transparent communication. The V storage facility is implemented in the kernel, and the system call interface is implemented in the kernel.

well as relatively easy to use, does not support application-level use of streaming.

These potential problems prompted a performance evaluation of our methods, with particular emphasis on the efficiency of file access. This emphasis on file access distinguishes our work from similar studies [10, 13]. The results of our study strongly support the idea of building a distributed system using distinct workstations connected by a high-speed local network to one or more file servers. Furthermore, we show that stream file access using the V kernel IPC facility is only slightly more expensive than a lower bound imposed by the basic cost of network communication. From this we conclude that relatively little improvement in performance can be achieved using protocols further specialized to file access.

2. V Kernel Interprocess Communication

The basic model provided by the V kernel is that of many small processes communicating by messages. A process is identified by a 32-bit globally unique process identifier or *pid*. Communication between processes is provided in the form of short fixed-length messages, each with an associated reply message, plus a data transfer operation for moving larger amounts of data between processes. In particular, all messages are a fixed 32 bytes in length.

The common communication scenario is as follows: A client process executes a *Send* to a server process which then completes execution of a *Receive* to receive the message and eventually executes a *Reply* to respond with a reply message back to the client. We refer to this sequence as a *message exchange*. The receiver may execute one of *MoveTo* or *MoveFrom* data transfer operations between the time the message is received and the time the reply message is sent.

The following sections describe the primitives defined in this paper. The *kernel* module is defined in the V kernel source code [1] for a complete description of the kernel facilities.

2.1. Primitives

Send(*message*, *pid*)

Send a message to the process identified by *pid*. The message is a 32-byte fixed-length message. The *message* parameter is a pointer to the message buffer. The *pid* parameter is a 32-bit process identifier. The *Send* operation is a blocking operation. The process calling *Send* will not return until the message has been received by the process identified by *pid*. The *Send* operation is a blocking operation. The process calling *Send* will not return until the message has been received by the process identified by *pid*.

Receive

(*pid*, *scope*) = *Receive*(*message*, *scope*, *reply*, *scope*)
Receive a message from the process identified by *pid*. The message is a 32-byte fixed-length message. The *message* parameter is a pointer to the message buffer. The *scope* parameter is a 32-bit process identifier. The *reply* parameter is a pointer to the reply message buffer. The *scope* parameter is a 32-bit process identifier. The *Receive* operation is a blocking operation. The process calling *Receive* will not return until the message has been received from the process identified by *pid*.

Reply(*message*, *pid*)

Send a 32-byte reply message to the process identified by *pid*. The message is a 32-byte fixed-length message. The *message* parameter is a pointer to the message buffer. The *pid* parameter is a 32-bit process identifier. The *Reply* operation is a blocking operation. The process calling *Reply* will not return until the message has been received by the process identified by *pid*.

ReplyWithSegment(*message*, *pid*, *segment*, *scope*, *scope*)

Send a reply message to the process identified by *pid* and transfer the data segment specified by *segment* and *scope* to the process identified by *scope*. The *message* parameter is a pointer to the message buffer. The *pid* parameter is a 32-bit process identifier. The *segment* parameter is a pointer to the data segment. The *scope* parameter is a 32-bit process identifier. The *ReplyWithSegment* operation is a blocking operation. The process calling *ReplyWithSegment* will not return until the message has been received by the process identified by *pid*.

MoveFrom(*segment*, *dest*, *src*, *scope*)

Copy data from the segment starting at *src* in the address space of *scope* to the segment starting at *dest* in the address space of *scope*. The *segment* parameter is a pointer to the data segment. The *dest* parameter is a pointer to the destination address. The *src* parameter is a pointer to the source address. The *scope* parameter is a 32-bit process identifier. The *MoveFrom* operation is a blocking operation. The process calling *MoveFrom* will not return until the data has been copied from the segment starting at *src* to the segment starting at *dest*.

MoveTo(*segment*, *dest*, *src*, *scope*)

Copy data from the segment starting at *src* in the address space of *scope* to the segment starting at *dest* in the address space of the process calling *MoveTo*. The *segment* parameter is a pointer to the data segment. The *dest* parameter is a pointer to the destination address. The *src* parameter is a pointer to the source address. The *scope* parameter is a 32-bit process identifier. The *MoveTo* operation is a blocking operation. The process calling *MoveTo* will not return until the data has been copied from the segment starting at *src* to the segment starting at *dest*.

GetFile(*segment*, *scope*)

Return a pointer to the segment identified by *segment* in the address space of *scope*. The *segment* parameter is a pointer to the data segment. The *scope* parameter is a 32-bit process identifier. The *GetFile* operation is a blocking operation. The process calling *GetFile* will not return until the segment has been located in the address space of *scope*.

pid = *GetFile*(*segment*, *scope*)

Return the process identifier associated with *segment* in the address space of *scope*. The *segment* parameter is a pointer to the data segment. The *scope* parameter is a 32-bit process identifier. The *pid* parameter is a 32-bit process identifier. The *pid* operation is a blocking operation. The process calling *pid* will not return until the process identifier has been located in the address space of *scope*.

2.2. Stream Access

The V kernel provides a mechanism for accessing files. The *File* module is defined in the V kernel source code [1] for a complete description of the kernel facilities. The *File* module provides a mechanism for accessing files. The *File* module is defined in the V kernel source code [1] for a complete description of the kernel facilities.

The *File* module provides a mechanism for accessing files. The *File* module is defined in the V kernel source code [1] for a complete description of the kernel facilities.

The *File* module provides a mechanism for accessing files. The *File* module is defined in the V kernel source code [1] for a complete description of the kernel facilities.

However, our experience has been that the V message primitives are easily and efficiently implemented over a local network. Moreover, we have found that the semantics of the primitives facilitated an efficient distributed implementation. The only major departure from Thoth was the explicit specification of segments in messages and the addition of the primitives *ReceiveWithSegment* and *SendWithSegment*. This extension was done for efficient page-level file access although it has proven useful under more general circumstances, eg. in pushing document state values to remote servers.

A second concern is the implementation of individual pay based systems. Before describing some of the implementation details of the individual systems, we first review aspects of the implementation that are common to the different systems of the bank.

2. Reserve operations are implemented directly in the bank instead of through a payment bank account. When the bank receives a payment order from a customer, it immediately debits the customer's account and credits the payee's account. This process is known as "direct debit" and is used for many types of payments, including utility bills, rent, and salaries.

100-443887-100

6. File page-level transactions require the original number of pages (i.e. two) because of the ability to append short segments to messages using *ReceiveWithSignature* and *SendWithSignature*.

In the current 2.0b Ethernet implementation, the top 4 bits of the logical host identifier are the physical network address of the controller, making the present identifier to network address mapping direct. In the 3.0b implementation, a 48-bit unique identifier is network address. When there is no other entry for the specified logical host, the average is broadcast. New "logical identifier" address" correspondence can be determined from average address. However, all such kind of dependence is absent in our model host identifier during broadcast.

[illegible]

process. With the exception of GSP's local operations with no process identifier parameters are implicitly local to the workstation.

GSP's use network broadcast to determine the mapping of a logical process identifier to real process identifier if the mapping is not known to the local kernel. Any kernel knowing the mapping can respond to the broadcast request. The addition of local and remote scopes was required to discriminate between server processes that serve only a single workstation and those that serve the network.

3.2. Remote Message Implementation

When a process identifier is specified to *Send* with a logical host identifier different from that of the local machine, the local pid validation test fails and *Send* calls *NonLocalSend* which handles transmission of the message over the network.

The *NonLocalSend* routine writes a *InterNet* packet on the network addressed to the host machine of the process or else broadcasts the packet if the host machine is not known. When the host containing the recipient process receives the packet, it creates an *alien* process descriptor to represent the remote sending process using a standard kernel process descriptor¹ and saves the message in the message buffer field of the *alien* process descriptor. When the receiving process replies to the message, the reply is transmitted back to the sender as well as being saved for a period of time in the *alien* descriptor. If the sender later receives a reply within the timeout period *T*, the original message is acknowledged by the sender's kernel. The receiving kernel then has the responsibility of removing messages by copying the message contents into the *alien* process descriptor with data represented by the *alien*. The kernel removes the acknowledged message by discarding the message and either acknowledging the reply message or else sending back a "reply pending" packet to the sending kernel if the reply has not yet been received. It also sends back a reply pending packet if it is found to contain a new message because an *alien* process descriptor can contain only one message because the *alien* process descriptor is unidirectional. The sending kernel schedules the sending process to receive the packet that the *alien* descriptor is a message to receive acknowledgment packet or a message to send which causes the *alien* reply message to be sent back to the sender.

The following sequence of events occur when the *alien* message is received by the kernel. The kernel first checks the *alien* process descriptor to see if it is a message to receive acknowledgment packet or a message to send which causes the *alien* reply message to be sent back to the sender.

3.3. Remote Message Reception

The *alien* process descriptor is a kernel process descriptor. It contains a pointer to the *alien* process descriptor which is a kernel process descriptor. It contains a pointer to the *alien* process descriptor which is a kernel process descriptor.

the correct sequence. *MoveTo* transmits the data to be moved in a sequence of sequentially-sized packets to the destination workstation and waits a single acknowledgement packet when all the data has been received. Given the observed low error rates of local networks, all communications on over bandwidth only a slight performance degradation. We have, however, implemented acknowledgement from the list currently received data packet in order to avoid the pitfall of repeated identical failures arising when back-to-back packets are consistently being dropped by the receiver. The implementation of *MoveTo* is similar except a *MoveTo* request is sent out and acknowledged by the requested data packet, essentially the reverse of *MoveTo*.

As in the local case, major operations arise with *MoveTo* and *MoveFrom* because, by their definition, there is no need for queuing or buffering of the data in the kernel. The *V* kernel copies the data directly from the process address space into the network interface, and directly from the network interface to the destination address space².

3.4. Remote Segment Access

The message and data transfer primitives provide efficient communication of small amounts and large amounts of data, but that is not the case of direct access of shared objects. However, page-level file access requires a large amount of data that is not efficiently handled by the *MoveTo* primitives when implemented over a local network.

V file access is implemented using an *I/O* protocol developed for *Vers* [4]. To read a page or block of a file, a client sends a message to the file server process specifying the file, block number, byte count and the address of the buffer into which the data is to be returned. The server reads the data off disk, checks to see the sender buffer using *MoveTo*, and then replies, confirming the amount of data read. In the common case of the data being less than a single network packet, the receiver's packet transmission is pending as substantially as for the first 2 for the first round and one for the reply. This is double the number of packets required for a sequential page-level file access protocol as used for *MoveTo* in *VS* [4] (page 400-401).

To access the portion, we make some modifications to the *alien* process descriptor. For, in order to provide the *alien* process descriptor and *InterNet* packet, we have to make a modification of *alien* to change its *alien* process descriptor to *alien* process descriptor, and the *alien* process descriptor to *alien* process descriptor. The *alien* process descriptor is a kernel process descriptor. It contains a pointer to the *alien* process descriptor which is a kernel process descriptor. It contains a pointer to the *alien* process descriptor which is a kernel process descriptor.

The *alien* process descriptor is a kernel process descriptor. It contains a pointer to the *alien* process descriptor which is a kernel process descriptor.

at least as long as a file block, a file write operation is reduced to a single two packet exchange. In this fashion, much more is sent across the link as single block data write operations are handled efficiently. The High-Performance operation also reduces the write packet in the case of a file read by combining the reply message packet with the data to be returned.

Advantages of this approach include:

1. The advantages of the Trueth FC model with the network performance of a VME with file access protocol.
2. Compatibility with existing host operation. The business systems may be retained exactly if in the future different users or if the scripter process operates a LIMA device.
3. Use of Asynchronous Transfer Mode (ATM) and High-Speed Data Transfer (HSDT) to provide security using ATM.

An expected objection to our solution is the apparent redundancy and complexity of the storage problem. We feel the redundancy may be overdone given the potential complexity of communication and synchronization between client and server programs. Also, it is not unexpected that there is some overlap between efficiency and elegance. Given that applications currently access system services through such interfaces that provide a procedural interface to the storage problem, it is not inappropriate to make some compromise in the storage level for efficiency.

We now turn to the discussion of the performance evaluation of the layout. We first define the term network penalty as a reasonably lower bound to the cost of network communication. Subsequently we discuss the efficiency of the layout, both in terms of storage penalty and efficiency.

4. Network Penalty

Our measurements of the VME are primarily relevant only to comparisons.

1. The cost of remote operation versus the cost of the corresponding local operation.
2. The cost of the entire data transfer versus the cost of the corresponding local operation.

An important issue is the comparison of the cost of the network communication to the cost of the local operation. The cost of the network communication is the cost of the data transfer plus the cost of the network communication. The cost of the local operation is the cost of the data transfer plus the cost of the local operation. The cost of the network communication is the cost of the data transfer plus the cost of the network communication. The cost of the local operation is the cost of the data transfer plus the cost of the local operation.

no error. The network penalty is a function of the processor, the network, the storage interface and the number of bytes transferred. It is the network data penalty for transporting the network between two systems, whether the data is stored or transferred. The network penalty is defined by assuming the data is stored in a buffer from the main memory of one computer to the main memory of another and the user not storing the data here for the equivalent to 1. The equivalent is a system of data for statistical purposes. The number are implemented at the data link layer under the network level to that as proposed or present including network operation in the results. The assumption of error-free transmission and low network utilization are good approximations of most local network environments.

Network penalty provides a more realistic minimum estimate than the data transfer time that is required by the physical network speed because it includes the processor and network interface delay. For instance, a 30 Mbit Ethernet can move 100 Mbit data can experience a penalty in 100 microseconds. However, the time to assemble the packet in the interface at the sending end and the time to transfer the packet out of the interface at the receiving end are comparable to the time for transmission. Thus, the time for the transfer from the point of view of the communicating software modules is at least two or three times as long as that implied by the physical transfer rate.

Measurement of network penalty was made using the equivalent 100 Mbit Ethernet. In all measurements, the network was operating at 100 Mbit/sec. The cost of the network communication is the cost of the data transfer plus the cost of the network communication. The cost of the local operation is the cost of the data transfer plus the cost of the local operation. The cost of the network communication is the cost of the data transfer plus the cost of the network communication. The cost of the local operation is the cost of the data transfer plus the cost of the local operation.

Network Penalty

Speed	Network Time	Network Penalty	Local Penalty
100	174	630	630
100	174	120	630
100	174	120	630
100	174	120	630
100	174	120	630

Table 1. Network Penalty (Network Time) (Network Penalty)

The network penalty is the cost of the network communication plus the cost of the local operation. The cost of the network communication is the cost of the data transfer plus the cost of the network communication. The cost of the local operation is the cost of the data transfer plus the cost of the local operation. The cost of the network communication is the cost of the data transfer plus the cost of the network communication. The cost of the local operation is the cost of the data transfer plus the cost of the local operation.

With our interface, the processor is required to copy the packet into the interface for transmission and one of its functions is reception (with the interface providing backpressure to packet buffering). From the copy that goes down, our design said that a DMA interface would significantly improve performance. However, we would prefer that a DMA interface never be used in favor of direct path access for two reasons. First, the kernel inspects a newly arrived packet as it enters the system from the network interface, allowing it to place parts of the packet into immediately to be sent locations. With a DMA interface, this copy would be unnecessary for the packet from DMAed into main memory. Similarly, an application that kernel dynamically constructs a network packet from disparate locations as it enters the system has the network interface. When DMA interfaces require the packet to be constructed in the contiguous area of memory, forcing the user to use complex copy operations. Finally, there is no network interface that we know of that a network DMA interface for the network that moves data faster than a 10 MB/s Network Interface Card is used in the SUN workstation. In general, the main benefit of a smart network interface appears to be in allowing the user processor rather than spending its operations for main use of network communication.

Our first job is to ensure that we have a strong and healthy economy and a strong and healthy environment. We have a lot of work to do in these areas, and we are committed to doing it. We will continue to work with the private sector and the public sector to create jobs and to protect our environment. We will also continue to work with the international community to address global challenges. We will continue to work with the people of the United States to ensure that they have the opportunity to live a better life. We will continue to work with the people of the world to ensure that they have the opportunity to live a better life. We will continue to work with the people of the United States and the people of the world to ensure that they have the opportunity to live a better life.

Using 1000 trials per equation and then selecting plus or minus 10 coefficients, our measurements should be enough to show 10 coefficients enough for the effect of variation in network load.

Table 5-1 gives the results of our measurements of average velocities and the number with the lowest reading on each machine using a 25% faster promoter and compared by a 50% slower promoter. Thus the difference is a third third operation. Indeed it gives the best statistical evidence of a third operation. The average third third and fourth give the same data as the first operation almost evenly and relatively. The difference between the two differences between the first and second operation. The third operation gives the correct point in the middle of the difference in part of the second operation. The third and fourth operation for the promoter time and for the operation on the two machines involved in the second operation of the operation. Table 5-2 gives the same measurements using a 25% faster promoter. The data for both promoters are given separately for the effect the promoter speed has on the third operation. As expected, the data for both operations, being dependent only on the promoter speed, are 25 percent faster on the 25 percent faster promoter. However, the table 25 percent improvement for same operation follows the promoter in the next significant improvement, being in our configuration and is not statistical evidence for the second, being at least on a slightly faster promoter.

A significant trend of increasing mortality along plant lateness gradients was found for *Myrica gale*. A fully significant interaction was observed between the gradient and the number of years since the last fire, with the greatest mortality occurring in the youngest stands. The interaction was not significant for the remaining species, and the mortality of the species was not related to the number of years since the last fire. The mortality gradient was not related to the number of years since the last fire for any of the species, and the mortality gradient was not related to the number of years since the last fire for any of the species.

Kernel Performance

Kernel Operation	Elapsed Time			Network Penalty	Processor Time	
	Local	Remote	Difference		Client	Server
GetTime	0.07	-	-	-	0.07	-
Send-Receive-Reply	1.00	1.18	2.18	1.00	1.79	2.30
MoveFrom: 1024 bytes	1.26	9.83	7.77	8.15	3.76	5.69
MoveTo: 1024 bytes	1.26	9.85	7.79	8.15	3.59	5.87

Table 5-1: Kernel Performance: 3 MB Ethernet and 80 MHz Processor (times in milliseconds)

Kernel Operation	Elapsed Time			Network Penalty	Processor Time	
	Local	Remote	Difference		Client	Server
GetTime	0.06	-	-	-	0.06	-
Send-Receive-Reply	0.77	2.54	1.77	1.30	1.44	1.79
MoveFrom: 1024 bytes	0.95	8.60	7.65	6.77	3.32	4.78
MoveTo: 1024 bytes	0.95	8.60	7.65	6.77	3.17	4.95

Table 5-2: Kernel Performance: 3 MB Ethernet and 10 MHz Processor (times in milliseconds)

view interprocess communication as transparent across machines when the speed ratio is so large. However, an alternative interpretation is to recognize that the remote operation adds a delay of less than 2 milliseconds, and that in many cases this delay is insignificant relative to the time necessary to generate a request in the server. Furthermore, the waiting at client workstation processor is tiny with the server load for only 1.04 milliseconds out of the total 2.14 milliseconds time taken by the 80 MHz processor). Then, one can afford the processor on one machine to, for instance, moving a server process to another machine if its request processing generally requires more than 0.07 milliseconds of processor time, i.e. the difference between the local time (Send-Receive-Reply time and the local processor time is 0.07 milliseconds).

5.4. Small-Program Traffic

The situation so far has been for a single pair of processes communicating over the network. Operating systems and network configurations should be able to support the network communication to applications with other problems. Some hardware is designed to handle multiple connections and some software is designed to handle the application to communication layer of the communication.

A pair of workstations connected by a network can be used to test the network communication. The network communication is tested by sending a message from one workstation to another workstation. The message is sent from the client workstation to the server workstation. The message is sent from the server workstation to the client workstation. The message is sent from the client workstation to the server workstation. The message is sent from the server workstation to the client workstation.

network in this fashion. Unfortunately, our measurements of this scenario turned up a hardware bug in our 3 MB Ethernet interface, a bug which caused many collisions to go undetected and show up as corrupted packets. Thus, the response time for the 80 MHz processor workstation in this case is 3.4 milliseconds. The increase in time from 1.18 milliseconds is accounted for almost entirely by the detected and retransmitted data (usually one per 200 packets) from the hardware bug. With standard network interfaces, we believe that the network can support any reasonable level of message communication without significant performance degradation.

A more critical resource is processor time. This is especially true for workstations which are servers that need to be the focus of a number of network connections. The problem just faced on workstations that a workstation is loaded to a point where 500 message exchanges per second, independent of the number of clients. The number is substantially lower for the server workstation, particularly when a remote client for the server processing is involved. The table summarizes an example in Table 5-1.

The network communication is tested by sending a message from one workstation to another workstation. The message is sent from the client workstation to the server workstation. The message is sent from the server workstation to the client workstation. The message is sent from the client workstation to the server workstation. The message is sent from the server workstation to the client workstation.

serving the client process we are measuring and otherwise idle. A later section considers multi-client load on the file server.

We first describe the performance of random page-level file access.

6.1. Page-Level File Access

Table 6-1 list the times for reading or writing a 512 byte block between two processes both local and remote using the 10 NFS processor. The times do not include time to fetch the data from disk but do indicate expected performance when data is buffered in memory. A page read involves the sequence of kernel operations: *Send-Receive-ReplyWithSegment*. A page write is *Send-ReceiveWithSegment-Reply*.

Random Page-Level Access

Operation	Elapsed Time			Network Penalty	Processor Time	
	Local	Remote	Difference		Client	Server
page read	1.31	5.96	4.25	1.89	2.50	3.28
page write	1.31	5.69	4.29	1.89	2.50	3.32

Table 6-1: Page-Level File Access: 512 byte pages (times in milliseconds)

The columns are to be interpreted according to the explanation given for similarly labeled columns of Tables 3-1 and 3-2. Thus, the time to read or write a page using these primitives is approximately 1.5 milliseconds more than the disk access latency for these operations.

There are several considerations that apply to the use of remote operations being faster than local operations. There are special cases of these described for simple message exchanges. First, the extra 4.2 milliseconds time for remote operations is relatively small compared to the time cost of the file system operations itself. In particular, disk access time can be estimated at 20 milliseconds (assuming 1000000 bps) and the time to transfer data at 2.5 milliseconds. The total time to read or write a 512 byte block is 22.5 milliseconds and a write time of 23.2 milliseconds, making the cost of the remote operations 20 percent faster than the local operations.

The second consideration is that the time to transfer data over the network is a significant factor. In this case, the server is a 10 NFS processor and the time to transfer data over the network is 1.89 milliseconds. This is a significant factor in the overall time to transfer data over the network.

Third, the time to transfer data over the network is a significant factor. In this case, the server is a 10 NFS processor and the time to transfer data over the network is 1.89 milliseconds. This is a significant factor in the overall time to transfer data over the network.

Fourth, the time to transfer data over the network is a significant factor. In this case, the server is a 10 NFS processor and the time to transfer data over the network is 1.89 milliseconds. This is a significant factor in the overall time to transfer data over the network.

Fifth, the time to transfer data over the network is a significant factor. In this case, the server is a 10 NFS processor and the time to transfer data over the network is 1.89 milliseconds. This is a significant factor in the overall time to transfer data over the network.

difference between the client processor time for remote page access and for local page access, namely 1.3 milliseconds. A processor cost of more than 1.3 milliseconds per request can be expected from the estimation made earlier using LOCUS figures.

These measurements indicate the performance when file reading and writing use explicit segment specification in the message and *ReceiveWithSegment* and *ReplyWithSegment*. However, a file write can also be performed in a more basic Thrift-like way using the *Send-Receive-MoveFrom-Reply* sequence. For a 512 byte write, this costs 8.1 milliseconds; file reading is similar using *MoveTo*. Thus, the segment mechanism saves 3.5 milliseconds on every page read and write operation, justifying this extension to the message primitives.

6.2. Sequential File Access

Sequential file access is the predominant pattern for file activity in most systems. Efficient file systems exploit this behavior to reduce the effect of disk latency by prefetching file pages (read-ahead) and aggressively flushing modified pages (write-behind). File systems and file transfer protocols typically implement mechanisms to reduce the effect of network latency on sequential file access.

Using the V kernel remote communication between a workstation and a file server, the file server can implement read-ahead and write-behind to reduce the effect of disk latency. However, there is no streaming in the network IFC to deal with network latency.

Two factors suggest that streaming can be done without a local cache on the server. First, local caches have a low hit ratio. Second, a local cache is a significant performance overhead. Third, a local cache is a significant performance overhead. Fourth, a local cache is a significant performance overhead. Fifth, a local cache is a significant performance overhead. Sixth, a local cache is a significant performance overhead. Seventh, a local cache is a significant performance overhead. Eighth, a local cache is a significant performance overhead. Ninth, a local cache is a significant performance overhead. Tenth, a local cache is a significant performance overhead.

The second factor is that the time to transfer data over the network is a significant factor. In this case, the server is a 10 NFS processor and the time to transfer data over the network is 1.89 milliseconds. This is a significant factor in the overall time to transfer data over the network.

many current unloading systems do, such program loading could achieve the same performance given in the table, independent of disk speed. Thus, we argue that *MoveTo* and *MoveFrom* with large transfer units provide an efficient program loading mechanism that is as fast as can be achieved with the given hardware.

7. File Server Issues

File server performance is a critical issue for diskless workstations. Unfortunately, we do not yet have experience with a V kernel-based file server. Thus, this section describes what we believe are the key issues and estimates performance without providing conclusive data. In general, we view the processor as the key resource to consider in file server performance issues, as argued earlier, the network bandwidth is plentiful and disk scheduling and buffering issues are identical to those encountered in conventional multi-user systems.

The number of workstations a file server can support can be estimated from processor requirements. If we estimate page read or write processing overhead as roughly 15 milliseconds for file system processing (from LOCUS) plus 13 milliseconds for bus operation (from Table 6-13), a page request costs about 7 milliseconds of processor time. Program loading appears to cost about 300 milliseconds for an average 64 kilabyte program. Estimating that 25 percent of the file requests are page requests, the average request costs 36 milliseconds. Thus, a file server based on the SUN workstation processor could support about 25 file requests a second. From this we estimate that one file server can serve about 10 workstations satisfactorily, but 20 or more active workstations would lead to excessive delays. However, a dedicated workstation cannot be easily be attached to multiple more workstations by adding more file server machines than the network would not mean as to a benchmark for file than 100 workstations.

For these programs, it is advantageous to view all of the donor program expenditures to conduct the programs under the same rules. Thus to fund the program has a significant and subsequently full income tax expenditure than to fund programs operating for a short time and during a lot of the time with operating out of the characteristics of the program and related activities with the IRS.

[illegible]

program, i.e. it is management oriented for performance.

B. Measurements with the 10 Mb Ethernet

Our limited access to a 10-Mb Ethernet has precluded basing our measurements on this standard local network. However, some preliminary figures using the 10-Mb Ethernet indicate the effect of using a faster network and slightly faster network interfaces. First, the remote message exchange time is 2.71 milliseconds using an 8-MHz processor, roughly the time for the 10-MHz processor on the 3-Mb network and .5 milliseconds better than the 8-MHz processor on the 3-Mb network. Second, the page read time is 5.72 milliseconds. Finally, the program loading time is much improved, achieving 255 milliseconds for a 64-kilobyte load using 16-KB transfer units. We have not identified to what degree the improvement is due to the faster network speed versus the differences in the network interfaces.

9. Related Work

There are a number of previous and concurrent efforts in providing communication mechanisms for distributed systems. For brevity, we compare our work with only a representative sample that encompasses the search for, and evaluation of, relevant models and implementations.

Speyer's remote reference study [13] considered the feasibility of implementing remote find and merge operations over a local network. Nakano's work on remote procedure calls [10] investigates network communication for procedure-based systems and argues as to what the V kernel provides for manage-based systems. Raskin and Seltmann's Agent kernel [12] implements a manage-agent with a number of features such as non-blocking message sending that are not provided by the V kernel. Finally, LOCKER [11] manages network communication using a UNIX-like kernel in the form of a network device file system.

One of the major factors of the success of Spitzer's and Hahn's work is using a super-regional form of communication (a place of knowledge and strategy were placed here for efficient performance). However, in place of Spitzer's stance, there is the idea of a super-regional strategy about money in a distributed system, which is especially the functionality provided by the processor. The V level provides a strong degree of separation between process and software provided provision of services in a software, multi-machine environment by setting a new level of abstraction in the form of a system.

Our research differs from others' primarily in our use of a highly sophisticated instrument that provides the actual thoughts and ideas of the standing difference in research. Furthermore, the VHS provides data on which to base a more precise and thorough by the addition of more research and research related. These data enabled us to see the differences between the actual long-term research and the research of previous years.

versus message-based systems, although it is not clear these differences result in any significant difference in overall performance.

The V kernel performance is roughly comparable to that of the software implementations developed by Spector and Nelson, allowing for the non-trivial differences in operation semantics and host processors. We would hypothesize that V kernel performance could be improved by a factor of 30 using microcode, similar to the improvement observed by Spector and Nelson for their primitives. Unfortunately, neither Spector nor Nelson provides results that afford a comparison with our file access results. In general, their work has concentrated on the speed of the basic mechanism and has not been extended to measure performance in a particular application setting.

In comparison to Accent, the V kernel provides a primitive form of message communication, and benefits accordingly in terms of speed, small code size and ability to run well on an inexpensive machine⁵ without disk or microcode support. For instance, Accent messages require an underlying transport protocol for reliable delivery because there is no client-level reply message associated with every *Send* as in the V kernel. We do not at this time have performance figures for Accent.

LOCUS does not attempt to provide applications with general network interprocess communication but exploits carefully honed problem-oriented protocols for efficient remote file access. It is difficult to compare the two systems from measurements available given the differences in network speeds, processor speeds and measurement techniques. However, from the specific comparisons with LOCUS presented earlier, we would expect overall file access performance for the V kernel to be comparable to LOCUS running on the same machines and network.

However, the memory requirements for the V kernel are about half that of LOCUS compiled for the PDP-11 and probably more like one fourth when LOCUS is compiled for a 32-bit processor like the 68000. Thus, for graphics workstations or process control applications, for instance, the V kernel would be more attractive because of its smaller size, real-time orientation and its provision of general interprocess communication. However, the V kernel does not provide all the functionality of the LOCUS kernel which includes that of the UNIX kernel and more. When required with V, these additional facilities must be provided by server processes executing either on client workstations or network server machines.

10. Conclusions

We conclude that it is feasible to build a distributed system using diskless workstations connected by a high-speed local network to one or more file servers using the V kernel IPC. In particular, the performance study shows that V kernel IPC provides satisfactory

performance despite its generality. Because the performance is so close to the lower bound given by the network penalty, there is relatively little room for improvement on the V IPC for the given hardware regardless of protocol and implementation used.

The efficiency of file access using the V IPC suggests that it can not only replace page-level file access protocols but also file transfer and remote terminal protocols, thereby reducing the number of protocols needed. We claim that V kernel IPC is adequate as a transport level for all our local network communication providing each machine runs the V kernel or at least handles the interkernel protocol. We do, however, see a place for these specific protocols in internetworking situations.

In addition to quantifying the elapsed time for various operations, our study points out the importance of considering processor requirements in the design of distributed systems. More experience and measurement of file server load and workstation file access behavior is required to decide whether file server processing is a significant problem in using diskless workstations.

The V kernel has been in use with the diskless SUN workstations, providing local and remote interprocess communication, since September 1982. It is currently 38 kilobytes including code, data and stack. The major use of the network interprocess communication is for accessing remote files. Our file servers are currently 6 VAX/UNIX systems running a kernel simulator and file server program which provides access to UNIX system services over the Ethernet using interkernel packets. A simple command interpreter program allows programs to be loaded and run on the workstations using these UNIX servers. Our experience with this software to date supports the conclusions of the performance study that we can indeed build our next generation of computing facilities [8] using diskless workstations and the V kernel.

Acknowledgements

We are indebted to all the members of the V research group at Stanford, which at this point includes two faculty members and roughly ten graduate students. In particular, we wish to thank Keith Lauts for his patient comments on a seemingly endless sequence of drafts and Tim Mann for his many contributions to the design and the implementation of the kernel. We would also like to thank the referees whose comments and suggestions helped to enhance the clarity of the paper.

References

1. F. Buxton, J.H. Howard and J.T. Manning. *Task Communication in DEDACS*. Proceedings of the 8th Symposium on Operating System Principles, ACM, November, 1977, pp. 23-31. Publication (Operating Systems Review 12(3)).
2. A. Bach, J. Burch, F. Buxton, V. Paxon. *The SUN Workstation Architecture*. Tech. Rep. 82, Computer Science Laboratory, Stanford University, March, 1982.

⁵ In addition to the workstation used for this study, the V kernel runs on a PDP-11.

3. D.R. Cheriton, M.A. Malcolm, I.S. Melen and G.R. Sagar. "Thoth, a Portable Real-time Operating System." *Comm. ACM* 22, 2 (February 1979), 105-115.

4. D.R. Cheriton. Distributed I/O using an Object-based Protocol. Tech. Rept. 81-1, Computer Science, University of British Columbia, 1981.

5. D.R. Cheriton. *The Thoth System: Multi-process Structuring and Portability*. American Elsevier, 1982.

6. D.R. Cheriton, T.P. Mann and W. Zwanevool. V-System: Kernel Manual. Computer Systems Laboratory, Stanford University.

7. Digital Equipment Corporation, Intel Corporation and Xerox Corporation. The Ethernet: A Local Area Network - Data Link Layer and Physical Layer Specifications, Version 1.0.

8. K.A. Lantz, D.R. Cheriton and W.I. Nowicki. Third Generation Graphics for Distributed Systems. Tech. Rept. STAN-CS-82-938, Department of Computer Science, Stanford University, February, 1983. To appear in *ACM Transactions on Graphics*.

9. R.M. Metcalfe and D.R. Boggs. "Ethernet: Distributed Packet Switching for Local Computer Networks." *Comm. ACM* 19, 7 (July 1976), 395-404. Also CSL-75-7, Xerox Palo Alto Research Center, reprinted in CSL-80-2.

10. B.J. Nelson. *Remote Procedure Call*. Ph.D. Th., Carnegie-Mellon U., 1981. published as CMU technical report CMU-CS-81-119.

11. G. Popok, B. Walker, J. Chow, D. Edwards, C. Kline, G. Rudisin, G. Thiel. LOCUS: A Network Transparent, High Reliability Distributed System. Proceedings of the 8th Symposium on Operating Systems Principles, ACM, December, 1981, pp. 169-177.

12. R. Rathid and G. Robertson. Accent: A Communication Oriented Network Operating System Kernel. Proceedings of the 8th Symposium on Operating Systems Principles, ACM, December, 1981, pp. 64-75.

13. A. Spitzer. "Performing Remote Operations Efficiently on a Local Computer Network." *Comm. ACM* 25, 4 (April 1982), 246-250.

14. D. Swinehart, G. McDowell and D. Boggs. WFS: A Simple Shared File System for a Distributed Environment. Proceedings of the 7th Symposium on Operating Systems Principles, ACM, December, 1979, pp. 9-23.



Accession For	
DTIC GRAFI	☒
DTIC TAB	☐
Unannounced	☐
Justification	☐
Distribution/Availability	
Availability Codes	
Avail and/or	Special

DATE
FILMED
8