

B.S. (12)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER	2. GOVT ACCESSION NO.	3. RECIPIENT'S CATALOG NUMBER
	AD-A088767	
4. TITLE (and Subtitle)		5. TYPE OF REPORT & PERIOD COVERED
(6) FLOW CONTROL AND ROUTING ALGORITHMS FOR DATA NETWORKS.		(9) Paper-Technical <u>papers</u>
7. AUTHOR(s)		6. PERFORMING ORG. REPORT NUMBER
(10) R. G./Gallager and S. J./Golestaani		(14) LIDS-P-1013
9. PERFORMING ORGANIZATION NAME AND ADDRESS		8. CONTRACT OR GRANT NUMBER(s)
Massachusetts Institute of Technology Laboratory for Information and Decision Systems Cambridge, Massachusetts 02139		ARPA Order No. 345/5-7-75 ONR/N00014-75-C-1183 ARPA Order-345
11. CONTROLLING OFFICE NAME AND ADDRESS		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS
Defense Advanced Research Projects Agency 1400 Wilson Boulevard Arlington, Virginia 22209		Program Code No. 5T10 ONR Identifying No. 049-383
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office)		12. REPORT DATE
Office of Naval Research Information Systems Program Code 437 Arlington, Virginia 22217		(11) July 1980
16. DISTRIBUTION STATEMENT (of this Report)		13. NUMBER OF PAGES
Approved for public release; distribution unlimited.		6 (12) 82
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		15. SECURITY CLASS. (of this report)
		Unclassified
18. SUPPLEMENTARY NOTES		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
19. KEY WORDS (Continue on reverse side if necessary and identify by block number)		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number)		
We consider flow control algorithms consisting of two parts: quasi-static flow control and dynamic flow control. The quasi-static part uses short term average information on network utilization to allocate maximum data rates and to determine routes for each user. The rates are allocated to achieve an optimal trade-off between assigned priority cost functions for each user and the cost of congestion in the network. This optimization can be done by a distributed algorithm and is essentially no more complicated than optimizing routing alone. The dynamic flow control has the function of admitting or rejecting individual units of traffic.		

DTIC
ELECTE
AUG 28 1980
S
A

420950

80 8 22 038

AD A088767

DDC FILE COPY

20. (Continued)

into the network so as to enforce the maximum allocated rates and to prevent congestion by smoothing out the fluctuations in buffer occupancy.

July 1980

LIDS-P-1013

To be presented at ICC'80, Atlanta Georgia, October 27-30, 1980.

Do not type in this block

ABSTRACT

We consider flow control algorithms consisting of two parts: quasi-static flow control and dynamic flow control. The quasi-static part uses short term average information on network utilization to allocate maximum data rates and to determine routes for each user. The rates are allocated to achieve an optimal trade-off between assigned priority cost functions for each user and the cost of congestion in the network. This optimization can be done by a distributed algorithm and is essentially no more complicated than optimizing routing alone. The dynamic flow control has the function of admitting or rejecting individual units of traffic into the network so as to enforce the maximum allocated rates and to prevent congestion by smoothing out the fluctuations in buffer occupancy.

I. INTRODUCTION

The sources of a data network can be qualitatively modeled as random processes with slowly varying statistics. Thus, for example, on a time frame of seconds and milliseconds, we would focus on individual message arrivals into the network, whereas on a time frame of minutes and seconds, we would focus on the changing message arrival rates and the initiations and terminations of user sessions. If the network is moderately to heavily loaded, then both random fluctuations and the changing statistics can cause the inputs to temporarily overload the message handling ability of the network. If the network has no control over the inputs, then the overloads will cause the network buffers to fill up; this causes inefficient utilization of the communication lines, the throughput drops rapidly, and the network is congested; even if the input rapidly drops back to normal levels, the network remains congested.

It should be clear from the above that a network with finite buffers and communication capacity requires some means of controlling the inputs (if only to throw away messages that cannot be buffered). Flow control, as used here, is the set of techniques used by the network to control the inputs. An excel-

lent review and bibliography of flow control techniques is given by Gerla and Kleinrock [1]. Our treatment here will ignore what [1] refers to as transport level flow control (i.e., control between user processes), since that is a separable issue. We also ignore buffer management strategies except to the extent that they exert control on the inputs. Finally we do not distinguish between internal congestion in the network and congestion caused by a destination that is slow in absorbing messages from the network. This latter form of congestion can be viewed as internal congestion by modelling the path from destination node to external destination as a low capacity network link.

The general approach of flow control strategies is quite clear; one wants to curtail the inputs when the buffers in the network are filling up. There are two difficulties in this apparently simple approach. The first is that instantaneous global knowledge about the network is not available at the nodes; providing even partial information about the state of the network requires resources otherwise available for increasing throughput. The second is the issue of fairness and priorities; if inputs must be curtailed, which inputs should be curtailed and by how much?

Accession	
File	
DDC	
Unannounced	
Justification	
By	
Distribution	
Availability	
Dist.	

A ~~20~~

We consider separating flow control into a quasi-static and a dynamic part. The quasi-static part determines allowable transmission rates for each session, considering priorities, fairness, and the expected level of congestion throughout the network. The dynamic part has the function of admitting or rejecting individual messages to the network so as to satisfy the above allowable rates and to smooth out the fluctuations in arrivals and buffer occupancies. Conventional flow control strategies are primarily dynamic according to this division, and their quasi-static part is generally limited to rejecting new sessions when the system is heavily loaded.

As an example, consider a network with end to end windowing where each session has a window of w packets. If there are m sessions each transmitting packets as fast as possible, then on the order of mw packets will be buffered in the network (temporarily ignoring acknowledgement times). As m increases, the network will eventually become congested. Conventional strategies can sometimes avoid this congestion by rejecting new sessions, but if m increases due to inactive sessions becoming more active, the buffers will eventually fill up. The problem is that as the loading increases, delay increases, and the rate of each session is decreased, but not quite enough to prevent congestion. What is needed in such a situation is a means of adapting the window sizes to varying network conditions. We show in section 3 how to use the quasi-static flow control developed in section 2 to adapt window sizes to optimal values (according to a flexible criterion of optimality).

II. QUASI-STATIC FLOW CONTROL AND ROUTING

Our objective in this section is to determine an optimal set of input rates for a network of given topology with given user demands. We take the viewpoint of a supplier of network resources who can lose revenue either by failing to meet user demands or by causing long delays and lost messages due to congestion. Thus for each session using the network, we create a cost function, $e(r)$, which is a decreasing function of the rate r allocated to the given session (see Figure 1). As r is decreased, the cost in user dissatisfaction clearly increases. We discuss the implications of different forms of this function later, but for the time being simply restrict it to be twice differentiable and convex (i.e., the second derivative is non-negative).

Similarly, for each communication link of the network, we create a cost function $g(f)$, (see Figure 2) which is an increasing function of the traffic flow F on that link. As the flow approaches the capacity of the link the ex-

pected queue length of messages waiting to traverse the link increases and the danger of congestion increases; $g(F)$ should represent the cost of this congestion danger. It seems appropriate to choose $g(F)$ to approach ∞ at the maximum link flow that can be handled, given the number of buffers in the node. Typically this maximum F will be somewhat smaller than the capacity of the link, and we can call it the effective capacity, C_e , of the link.

A reasonable choice for $g(F)$, then, in keeping with the Kleinrock independence assumption for network queueing, is

$$g(F) = \frac{F}{C_e - F} \quad (1)$$

Our objective, now, is to form an aggregate cost function as the sum of the individual session cost functions and link cost functions. We then choose the session rates and network routes to minimize the aggregate cost function. In the remainder of this section, we set up the minimization problem precisely, give its solution, and then discuss the choice of session cost functions. Section three shows how the optimization can be implemented in a distributed way and relates the quasi-static solution to dynamic algorithms.

Suppose that the network has N nodes (switches) denoted by the integers $1, \dots, N$, and for notational simplicity, assume that for each pair (i, j) of nodes, there is at most one user session originating at node i and destined for node j . Let r_{ij}^d be the rate (in bits per second) desired by the user for this session and let r_{ij} be the rate allocated by the network.

Thus

$$0 \leq r_{ij} \leq r_{ij}^d \quad (2)$$

If there is no session from i to j , or if the session is currently inactive, we take $r_{ij}^d = 0$.

Let the network have L directed links denoted $1, \dots, L$. Each link goes out from some node i into some node j . Let $O(i)$ be the set of links going out from node i , and let $I(i)$ be the set of links going into node i . A full duplex communication channel between nodes i and j is regarded as two links, one going out from i and into j and the other going out from j and into i .

Let F_l be the user traffic flow (in bits/sec.) on link l and let f_{lj} be the part of that traffic belonging to sessions with destination node j . We then have the constraint equations:

$$\sum_{j=1}^N f_{lj} = F_l ; 1 \leq l \leq L \quad (3)$$

$$f_{lj} \geq 0 ; 1 \leq l \leq L, 1 \leq j \leq N \quad (4)$$

$$r_{ij} + \sum_{l \in I(i)} f_{lj} = \sum_{l \in O(i)} f_{lj} ;$$

$$1 \leq i, j \leq N; i \neq j \quad (5)$$

Equation (3) states that the total user traffic is simply the sum of the traffics to all of the destinations. Equation (5) states that the total traffic coming into node i destined for node j (the exogenous traffic r_{ij} , plus the traffic passing through i from other nodes) is equal to the total traffic going out. If f_{lj} was taken to be a short term average, then (5) would not be quite correct, and in fact the left side minus the right side would represent the rate of buffer buildup at node i for traffic going to j . Our viewpoint here is somewhat different, and is in fact the essence of the quasi-static approach. We are allocating session rates, and choosing flows which would be feasible over the long term with stationary inputs at the allocated rates. Any set of flows satisfying (3)-(5) can be achieved for the given input rates by a routing policy which, for each node i and destination j , routes the traffic to j out of i on outgoing links proportionally to the terms on the right hand side of (5); conversely the long term flows resulting from any routing policy (given that data is not thrown away inside the network) will satisfy (3)-(5).

The objective function that we want to minimize is given by

$$J(\vec{F}, \vec{r}) = \sum_{l=1}^L g_l(F_l) + \sum_{i=1}^N \sum_{j=1, j \neq i}^N e_{ij}(r_{ij}) \quad (6)$$

where $g_l(F_l)$ is the cost of congestion with flow F_l on link l and $e_{ij}(r_{ij})$ is the cost of restricting the session i, j to the rate r_{ij} .

Analytically, our problem is to minimize $J(\vec{F}, \vec{r})$ subject to the linear constraints (2)-(5). Applying Kuhn-Tucker theory, we are led to the following necessary and sufficient conditions for minimizing J . [2,3].

Theorem 1: Assume that g_l and e_{ij} have first and second derivatives satisfying $g'_l > 0$,

$g''_l > 0$, $e'_{ij} < 0$, $e''_{ij} > 0$ for $1 \leq l \leq L$, $1 \leq i, j \leq N$, $i \neq j$. Necessary and sufficient conditions on $\vec{F}^*$, $\vec{r}^*$ to minimize (6) subject to (2)-(5) are that a set of numbers λ_{ij} , $1 \leq i, j \leq N$ exist with $\lambda_{ii} = 0$, $1 \leq i \leq N$, such that a) for all i, j, k ,

$$g'_l(F_l^*) + \lambda_{kj} \geq \lambda_{ij} \quad (7)$$

for $l \in I(k)$, $l \in O(i)$, with equality if $f_{lj}^* > 0$.

and b), for all i, j such that $r_{ij}^d > 0$,

$$e'_{ij}(r_{ij}^*) + \lambda_{ij} \begin{cases} = 0 & \text{for } 0 < r_{ij}^* < r_{ij}^d \\ \leq 0 & \text{for } r_{ij}^* = r_{ij}^d \\ \geq 0 & \text{for } r_{ij}^* = 0 \end{cases} \quad (8)$$

If multiple sessions go from i to j , then r_{ij} in constraint (5) should be replaced by the sum of the allocations for i, j sessions, and (8) must be satisfied individually for each session.

We can interpret $g'_l(F_l^*)$ as the incremental cost of traffic on link l . If there is a path carrying traffic with destination j from node i to node j , then λ_{ij} is the incremental cost of traffic on that path. Equation (7) then states that all traffic flows on paths of minimum incremental cost. If one views $g'_l(F_l^*)$ as the "length" of link l , then (7) states that all traffic flows by shortest routes. This condition is well known when the rates are fixed and one is optimizing only over routing.

We will call $-e'_{ij}(r_{ij})$ the priority function for session i, j . It is the incremental gain for additional allocation to i, j . Equation (8) states that r_{ij}^* is set between 0 and r_{ij}^d to make this incremental gain as close as possible to λ_{ij} (i.e., equal to λ_{ij} except at the end points).

Note that the optimality conditions (7) and (8) depend only on the incremental link costs g'_l and the priority functions $-e'_{ij}$. This means that arbitrary constants could be added to g_l or e_{ij} without affecting the optimum.

Note also that the optimum point is independent of r_{ij}^d over the range of $r_{ij}^d > r_{ij}$. This is a desirable feature for a flow control

strategy, preventing users who are being actively flow controlled from attempting to increase their share of the resources by increasing their demands.

The maximum average network loading permitted by the flow control can be adjusted by scaling the priority functions. This scaling factor should depend, of course, on how effective the dynamic component of the flow control is in preventing congestion for a given short term average load level.

Some insight into the choice of priority functions can be obtained by considering the priorities of the form

$$-e'_{ij}(r_{ij}) = \left(\frac{a_{ij}}{r_{ij}} \right)^{b_{ij}} \quad (9)$$

The factor a_{ij} in (9) should be taken as a measure of the typical rate required for the session. For example, two sessions between nodes i and j with a given value of a_{ij} are equivalent to a single session with twice the above value for a_{ij} . The factor b_{ij} is more a measure of the importance of the session. The larger b_{ij} is, the less the session will be cut back with increasing network loading. Naturally, if b_{ij} is large for all sessions, this effect disappears, since the incremental costs simply increase with increasing demands until the allocated rates are feasible.

In considering algorithms to solve for the optimum rates and routes of (7) and (8), we start with a very general approach that reduces the problem to a quasi-static routing problem. For each session i, j , consider a fictitious communication link l' that carries all of the traffic, $r_{ij}^d - r_{ij}$, that is offered by the user but rejected by the flow control. Letting $O'(i)$ be the outgoing links from node i including these fictitious links, the flow constraint equation (5) becomes:

$$r_{ij}^d + \sum_{l \in I(i)} f_{lj} = \sum_{l \in O'(i)} f_{lj} \quad (5')$$

Now consider imposing a cost function on the fictitious link l' that carries the session i, j rejected traffic (see fig. 3):

$$g_{l'}(F_{l'}) = e_{ij}(r_{ij}^d - F_{l'}) \quad (10)$$

Since $r_{ij}^d - F_{l'}$ is just the accepted traffic for session i, j , $g_{l'}(F_{l'})$ is the same as

$e_{ij}(r_{ij})$. This means that the cost function (7) is equal to $\sum_l g_l(F_l)$ when the sum is taken over real and fictitious links.

The minimization of this cost function subject to (2)-(5') is now just a conventional quasi-static routing problem. In simple terms, we have replaced the question of how much traffic to allocate to session i, j with the equivalent question of how to apportion the traffic between the real and fictitious links. There are a variety of distributed algorithms for optimizing quasi-static routing in the literature [4]-[7], and any of them can now be used to jointly optimize quasi-static flow control and routing.

III. DYNAMIC QUASI-STATIC FLOW CONTROL

The major difficulty with considering flow control in terms of fictitious links in a routing problem is that this approach is not easily combined with dynamic flow control. To illustrate both this problem and possible alternative approaches, consider the use of end to end windowing for dynamic flow control [1], [8]. In this situation, the allocated traffic rates are determined by the window sizes. Thus the quasi-static part of the flow control is to choose the window sizes that lead to the optimum rate allocation. Unfortunately, if a single window size is changed in a network, this typically changes delays throughout the network, and thus changes the allocated rates for all the sessions.

We can find the change in rates for a given change in window sizes by ignoring flow control for the moment and finding the relation between the expected number of packets outstanding in the network for each session and the rate of each session, assuming constant routes. Number the sessions from 1 to m , letting r_m be the rate of session m and q_{ml} be the fraction of session m to use link l (note that for unifilar routes, q_{ml} takes on only 0 or 1 values). Let $t_l(F_l)$ be the average delay per packet on link l , including processing and queueing delays at the input to link l . Let $\bar{l}$ be the average packet length (where by packet we mean the transmission unit in terms of which windows are defined). Then the average number of packets for session m at link l is given by $r_m q_{ml} t_l(F_l) / \bar{l}$. This quantity, summed over l , is the average number of packets traversing the network at a given instant. Let α_m be the average delay for acknowledgements for session m . We assume priority for acknowledgement traffic so that α_m is independent of network delays. The averaged number of unacknowledged packets of session m out in

the network is then

$$w_m = \frac{r_m}{\Gamma} \left[\sum_l q_{ml} t_l(F_l) + \alpha_m \right] \quad (11)$$

Substituting $F_l = \sum_m r_m q_{ml}$ into (11) and differentiating, we get the matrix equation

$$P = T + RQMQ^T \quad (12)$$

where P_{mk} is $\Gamma \partial w_m / \partial r_k$, T is a diagonal matrix with $T_{mm} = \sum_l q_{ml} t_l(F_l) + \alpha_m$, R is a diagonal matrix with $R_{mm} = r_m$, Q is a matrix with elements q_{ml} , and M is a diagonal matrix with $M_{ll} = \partial t_l(F_l) / \partial F_l$.

Equation (11) relates window sizes for sessions to rates under the assumption that each user transmits as fast as possible subject to the window limitation. Golestaani [2] shows that the matrix P is invertible, which means that incremental changes in window sizes lead to uniquely defined changes in session rates. He shows, moreover, that (11) uniquely specifies the rates in terms of the window sizes for given increasing functions $t_l(F_l)$, given routing Q , and given α_m and Γ . Finally, he demonstrates a counter example to this uniqueness if acknowledgements are not sent at high priority.

Now consider the objective function (6) expressed as a function of the session rates r and the routing Q , $J^*(r, Q)$. It is easy to calculate $\partial J^* / \partial r_m$ as part of distributed routing algorithms [4]; it is simply the incremental cost of traffic on the routes for the session minus the priority function for the session. It is reasonable in a joint flow control and routing algorithm to increase (decrease) the window size for session m when $\partial J^* / \partial r_m$ is negative (positive). Assume in fact that one changes the window size of each session by an incremental amount $\Delta w_m = -\eta r_m \partial J^* / \partial r_m$. To first order, the change in J^* is then

$$\Delta J^* = \sum_m -\frac{\partial J^*}{\partial r_m} \Delta w_m \quad (13)$$

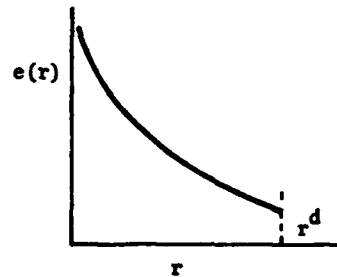
Using (12), this becomes

$$\Delta J^* = -\frac{\eta}{\Gamma} (V \cdot J^*) [TR + RQMQ^T R] (V \cdot J^*)^T$$

where $V \cdot J^*$ is the gradient of J^* with respect to w . Since T , M and R are positive diagonal matrices, it is easy to see that ΔJ^* is neg-

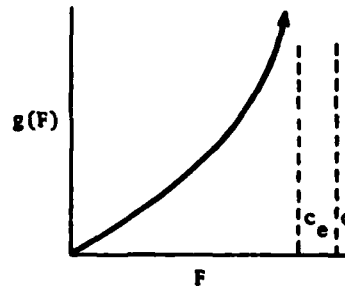
ative if $J^* \neq 0$, and this approach leads to a descent algorithm.

In the above, we have shown that it is possible to vary window sizes in accordance with a distributed routing algorithm so as to reduce the quasi-static objective function. For simplicity, a number of details have been omitted, such as the maintenance of positive session rates, the need for an artifact to achieve incremental variations in integer window sizes, and the combined effect of routing and window variations. Our objective has not been to demonstrate a finished algorithm, but rather to point the way to a wide class of algorithms that combine a quasi-static method of allocating rates with a dynamic method of smoothing traffic.



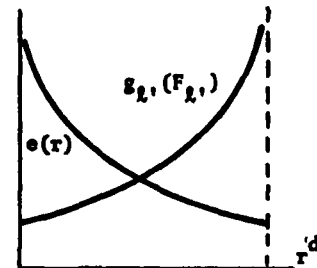
Cost $e(r)$ of Rate Limitation r on a User

Figure 1



Congestion Cost of Link Flow F

Figure 2



Relation of Session Cost $e(r)$ to Fictitious Link Cost

Figure 3

R. G. Gallager* and S. J. Golestaani**

*Mass. Inst. of Tech., Cambridge, Ma. 02139 USA; **Isfahan Un. of Tech., Isfahan, Iran

REFERENCES

- [1] M. Gerla and L. Kleinrock, "Flow Control: A Comparative Study", IEEE Trans. Commun., Vol. COM-28, pp. 553-574, April 1980.
- [2] S. J. Golestaani, "A Unified Theory of Flow Control and Routing in Data Communication Networks", Report LIDS-TH-963, January 1980, Laboratory for Information and Decision Systems, Mass. Inst. of Tech., Cambridge, MA 02139.
- [3] D. Luenberger, Introduction to Linear and Non Linear Programming, Addison-Wesley, 1973, Chapter 10.
- [4] R. G. Gallager, "A Minimum Delay Routing Algorithm Using Distributed Computation", IEEE Trans. Commun., Vol. COM-25, pp. 73-85, January 1977.
- [5] D. P. Bertsekas, "A Class of Optimal Routing Algorithms for Communication Networks", Proceedings of Fifth International Conference on Computer Communications (ICCC-80), November 1980.
- [6] A. Segall, "Optimal Distributed Routing for Virtual Line-Switched Data Networks", IEEE Trans. Commun., Vol. COM-27, pp. 201-209, January 1979.
- [7] T. E. Stern, "A Class of Decentralized Routing Algorithms Using Relaxation", IEEE Trans. Commun., Vol. COM-25, pp. 1092-1102, October 1977.
- [8] R. Kahn and W. Crowther, "Flow Control in a Resource Sharing Computer Network", IEEE Trans. Commun., Vol. COM-20, pp. 539-546, June 1972.

ACKNOWLEDGEMENT

This research was conducted at the M.I.T., Laboratory for Information and Decision Systems with partial support provided by grant DARPA/ONR-N00014-75-C-118e and grant NSF/ECS-7919880.

Robert G. Gallager received the S.B. degree in Electrical Engineering from the University of Pennsylvania in 1953 and the S.M. and Sc.D. degrees in Electrical Engineering from the Massachusetts Institute of Technology in 1957 and 1960 respectively. He has been at the Massachusetts Institute of Technology since 1956 and is currently Professor of Electrical Engineering and Computer Science and Associate Director of the Laboratory for Information and Decision Systems. He is also a consultant to

Codex Corporation. Professor Gallager is the author of the text book Information Theory and Reliable Communication (Wiley and Sons, 1968), and was awarded the IEEE Baker Prize Paper Award in 1966 for the paper "A Simple Derivation of the Coding Theorem and Some Applications". He is a member of the Board of Governors of the IEEE group on Information Theory, and was Chairman of the group in 1971. He is a Fellow of the IEEE and a member of the National Academy of Engineering. His major research interests are data communication networks, information theory, and communication engineering.

Seyyed J. Golestaani received the B.S. degree in Electrical Engineering from Arya-Mehr University of Technology in 1973 and the S.M. and Ph.D. degrees in Electrical Engineering and Computer Science from the Massachusetts Institute of Technology in 1976 and 1980 respectively. While studying at M.I.T., he was a research assistant in the Data Network Group, working in the areas of flow control, routing, and line protocols for error detection and retransmission. Since February 1980, he has been on the Electrical Engineering faculty at Isfahan University of Technology, Isfahan, Iran.

Distribution List

Defense Documentation Center Cameron Station Alexandria, Virginia 22314	12 Copies
Assistant Chief for Technology Office of Naval Research, Code 200 Arlington, Virginia 22217	1 Copy
Office of Naval Research Information Systems Program Code 437 Arlington, Virginia 22217	2 Copies
Office of Naval Research Branch Office, Boston 495 Summer Street Boston, Massachusetts 02210	1 Copy
Office of Naval Research Branch Office, Chicago 536 South Clark Street Chicago, Illinois 60605	1 Copy
Office of Naval Research Branch Office, Pasadena 1030 East Greet Street Pasadena, California 91106	1 Copy
New York Area Office (ONR) 715 Broadway - 5th Floor New York, New York 10003	1 Copy
Naval Research Laboratory Technical Information Division, Code 2627 Washington, D.C. 20375	6 Copies
Dr. A. L. Slafkosky Scientific Advisor Commandant of the Marine Corps (Code RD-1) Washington, D.C. 20380	1 Copy

Office of Naval Research
Code 455
Arlington, Virginia 22217

1 Copy

Office of Naval Research
Code 458
Arlington, Virginia 22217

1 Copy

Naval Electronics Laboratory Center
Advanced Software Technology Division
Code 5200
San Diego, California 92152

1 Copy

Mr. E. H. Gleissner
Naval Ship Research & Development Center
Computation and Mathematics Department
Bethesda, Maryland 20084

1 Copy

Captain Grace M. Hopper
NAICOM/MIS Planning Branch (OP-916D)
Office of Chief of Naval Operations
Washington, D.C. 20350

1 Copy

Advanced Research Projects Agency
Information Processing Techniques
1400 Wilson Boulevard
Arlington, Virginia 22209

1 Copy

Dr. Stuart L. Brodsky
Office of Naval Research
Code 432
Arlington, Virginia 22217

1 Copy

Captain Richard L. Martin, USN
Commanding Officer
USS Francis Marion (LPA-249)
FPO New York 09501

1 Copy