

AD-A157 739

PERFORMANCE TRADEOFFS AND HIERARCHICAL DESIGNS OF
DISTRIBUTED PACKET-SWIT. (U) CALIFORNIA UNIV LOS
ANGELES SCHOOL OF ENGINEERING AND APPLIED..

1/2

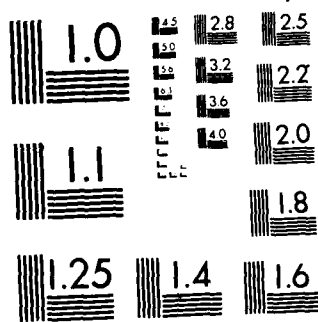
UNCLASSIFIED

G AKAUIA ET AL. SEP 79 UCLA-ENG-7952

F/G 17/2

NL





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

AD-A157 739

UCLA-ENG-7952
SEPTEMBER 1979

ARPA CONTRACT NO. MDA 903-77-C-0272

PERFORMANCE TRADEOFFS AND HIERARCHICAL DESIGNS OF
DISTRIBUTED PACKET-SWITCHING COMMUNICATION NETWORKS

G. AKAVIA and L. KLEINROCK

COMPUTER SYSTEMS
MODELING AND ANALYSIS GROUP

APPROVED FOR PUBLIC RELEASE;
DISTRIBUTION IS UNLIMITED (A)

COMPUTER SCIENCE DEPARTMENT

School of Engineering and Applied Science
University of California
Los Angeles



Serial Number 350,920

Filing Date 22 Feb 1981

Inventor Wolden B. Whittaker

(1)

NOTICE

The Government-owned invention described herein is available for licensing. Inquires and requests for licensing information should be addressed to:

DEPARTMENT OF THE NAVY
Chief of Naval Research (Code 302)
Arlington, Virginia 22217

3/18/75

UCLA-ENG-7952
SEPTEMBER 1979

**Performance Tradeoffs and Hierarchical Designs
of Distributed Packet-Switching Communication Networks**

by

Gideon Akavia
Computer Science Department
Technion, Israel Institute of Technology
Haifa, Israel

Leonard Kleinrock
Computer Science Department
University of California
Los Angeles

**Computer Science Department
School of Engineering and Applied Science
University of California
Los Angeles**

A-1

This research is an extension of the Ph.D. dissertation by Gideon Akavia under the Chairmanship of Leonard Kleinrock and was sponsored by the Advanced Research Projects Agency of the Department of Defense under contract MDA 903-77-C-0272.

Abstract

→ The design
We consider a distributed communication system with many terminals wishing to communicate with each other. When the terminals are distributed in space we must face the following questions: What scheme can control the access to the communication resources in an effective way? What tradeoffs are basic to the design of such a communication system? What is the role of hierarchies in organizing large communication nets? How should a large network be decomposed into smaller parts? What cost versus performance gains can be achieved by such a decomposition?

In attacking these questions we consider two technologies - line and broadcast - and two kinds of systems - *centralized systems*, in which messages originate in the distributed terminals but are directed to one common destination, and *networks*, in which both sources and destinations of messages are distributed.

We assume that the traffic to be carried and the necessary performance are specified and that the goal is to minimize the necessary cost. We define *quality* and *burstiness* and find the following: Dedicating channels is reasonable when the traffic is steady (i.e., not bursty), but when the traffic is bursty the cost of simple dedicated-channel systems grows too fast with the number of terminals. ALOHA is good when the traffic is bursty, but bad when the traffic is steady. Neither ALOHA nor dedicated channels are good when the traffic is of medium burstiness.

When given a broadcast channel, choosing the transmission range involves the following tradeoff: A long range enables messages to reach their destinations in a few hops, but increases the amount of traffic competing for the channel at every point.

In the first paper we calculate optimal transmission range. When choosing this optimal range, ALOHA networks gain a self adjusting capability, which makes heavily loaded ALOHA networks far better than centralized ALOHA systems. It is therefore harder to improve ALOHA networks than ALOHA centralized systems; power groups lead to a smaller relative improvement, while a hierarchy of ALOHA levels, with only a small population contending at the top level, can improve centralized systems but does not improve networks.

In the second paper we show that by introducing regular hierarchical structures the cost of bursty systems can be significantly reduced, and that the optimal structure must be balanced. In line systems the improvement follows from shortening individual lines, while in broadcast systems the improvement follows from spatial reuse.

The cost of the best bursty line system grows with the dimensionality of the space in which terminals are distributed. The cost of the best bursty broadcast system is similar to the cost of one-dimensional line systems and is independent of dimensionality. It follows that bursty broadcast systems have an advantage over line systems in two or more dimensions

PREVIOUS PAGE
IS BLANK

Organizing a two-dimensional network imposes a tessalation on the plane. When using the best number of levels, as a function of burstiness, tessalating the plane with hexagonal tiles (and forming a triangular network of communication lines) is usually optimal.

✓ In the third paper we show that mixed-mode systems, using ALOHA in a bottom level and dedicated channels in a top level, can be very good for medium burstiness since they can trade the amount of interference in the random access level against the number of dedicated channels in the top level. By choosing the right mix, such networks can become insensitive to the limitations of both access schemes.

CONTENTS

	Page
Abstract.....	iii
On a Self Adjusting Capability of ALOHA Networks	
Abstract.....	1
1. Introduction.....	1
2. Adjusting The Transmission Range.....	3
3. Using a Fixed Range.....	11
4. Capture, Power Groups and Partial Coordination.....	13
5. Multi-Level ALOHA.....	19
6. Conclusions.....	24
References.....	23
Hierarchical Use of Dedicated Channels	
Abstract.....	27
1. Introduction.....	27
2. Designing Distributed Communication Systems.....	30
3. Decomposition and Resource Allocation.....	32
4. Regular Hierarchical Structures.....	34
5. A Lower Bound?.....	37
6. Dedicated Broadcast Channels.....	38
7. Hierarchical Organization of Non-Bursty Line Systems.....	40
8. Distributed Dedicated-Line Networks.....	47
9. Hierarchical Line Networks.....	51
10. The Geometry of Networks.....	52
11. Conclusions.....	60
Appendix.....	61
References.....	62
On the Advantage of Mixing ALOHA and Dedicated Channels	
Abstract.....	65
1. Introduction.....	65
2. The 'No Interference' Case.....	66
3. The 'Full Interference' Case.....	68
4. Interacting ALOHA Subsystems.....	70
5. Sharing or Splitting?.....	74
6. Mixing with a General Random-Access Scheme.....	77
7. Are Three Levels Ever Necessary?.....	79
8. Two-Level Mixed-Mode Networks.....	83
9. Improving the Random Access Part.....	91
10. More than One Dedicated Level?.....	93
11. Conclusions.....	97
References.....	97
Computer Systems Modeling and Analysis Group Report Series	101

On a Self Adjusting Capability of ALOHA Networks

Abstract

We consider a distributed communication network with many terminals which are distributed in space and wishing to communicate with each other using a common radio channel. Choosing the transmission range in such a network involves the following tradeoff: a long range enables messages to reach their destinations in a few hops, but increases the amount of traffic competing for the channel at every point.

With the help of a simple model we analyze this tradeoff for ALOHA networks, and give the optimal range. When choosing this optimal range, as a function of specified traffic and delay parameters, ALOHA networks demonstrate an important self adjusting capability. This capability to adjust to traffic makes heavily loaded ALOHA networks far better than *centralized* ALOHA systems (in which all messages must reach one common destination).

Dividing a terminal population into *power groups* can improve any ALOHA system, especially when the traffic is split between groups in an appropriate way, which we demonstrate. But since ALOHA networks are hurt by destructive interference less than centralized ALOHA systems it is harder to improve them. Using power groups can significantly improve centralized systems, but will lead to a smaller relative improvement in ALOHA networks.

Decomposing the system into a *hierarchy* of ALOHA levels, with only a small population contending at the top level, can improve centralized systems but does not improve networks.

1. Introduction

Consider a large number of terminals, physically distributed over a large geographic region. If all terminals wish to communicate with one destination we shall call the system *centralized* and the common destination the station. Assuming the communication resource available is a radio channel of a given bandwidth, how should this common channel be shared among the terminals?

If the terminals were co-located in the same place, the best way to use the channel is to form a queue of busy terminals (i.e., those having anything to transmit) and to let them use the full bandwidth available one after the other. Forming one queue is much better than giving each terminal a fraction of the bandwidth, and letting each terminal queue its own messages [1].

It is no trivial matter to have all terminals form one queue when the terminals are numerous and distributed over large distances. Of special interest, then, is the ALOHA approach, which invests no resources in coordination and control of terminals. When using the (unslotted) ALOHA scheme each terminal transmits whenever it has a message ready. If more than one terminal is transmitting at the same time a conflict will occur in the use of the radio channel, and we shall assume at first that all messages involved in such a collision will be destroyed. When the destruction of its message becomes known to the terminal it will, after a somewhat randomized delay, retransmit the message. We shall not specify how the failure of its message becomes known to the terminal, but assume that this knowledge is free.



Schemes based on the ALOHA idea have been extensively treated [2,3,4]. ALOHA is obviously good when the system is lightly utilized and destructive interference is not very likely. When the load is heavy a significant fraction of the transmissions will fail as a result of collisions.

The wasteful effect of collisions can be reduced if all transmissions are of the same length [5]. This is usually achieved by breaking long messages into packets of a fixed maximum size. In this paper, we assume that this is always done and despite the fact that one message may result in several packets we assume that arrival of separate packets into our system is independent, and that the total arrival process is Poisson.

The wasteful effect of collisions can be further reduced if time is slotted (where each slot has a duration which is equal to a packet transmission time) and if terminals are constrained to start transmitting only at the beginning of a slot. The resulting access scheme is called Slotted ALOHA, and the maximum fraction of the time slots it can use for successful transmissions is known to be $1/e$ [6].

Let us choose the data unit so that the average length of a message is equal to 1. This is simply a convenient normalization, which is equivalent to measuring communication capacity in messages (of an average length) per second, instead of measuring in bits per second. The throughput-delay performance of the ALOHA schemes is not described by a simple analytic expression [3]. For simplicity we shall use the following ad-hoc expression to describe the performance of the ALOHA schemes

$$T = \frac{1}{C - eS} \quad (1)$$

Here T is the average response time of the system and S is the system throughput (messages per slot). We shall assume that this expression describes the optimum envelope of slotted ALOHA and unslotted ALOHA performance curves. (For $S \rightarrow 0$ it describes unslotted ALOHA, for $S/C \rightarrow 1/e$ it describes slotted ALOHA.) Equation (1) is a simple two-parameter approximation, that reproduces the known behavior when $S=0$ and when $S/C=1/e$. For a similar three-parameter approximation see [15].

Assume that the throughput S and the acceptable delay T are specified, and that we seek an access scheme that will minimize the necessary system capacity C . For most purposes it is sufficient to specify the communication needs by the dimensionless product ST , whose inverse we shall call *burstiness* [7,16,17]. We shall define the *quality* [7] of an arbitrary access scheme as the inverse ratio between the capacity necessary when using this scheme and the capacity necessary when using the *best possible scheme*, in which messages form one queue and share one channel. When messages arrive independently and their lengths are exponentially distributed the best scheme is the M/M/1 queue, in which we have $C = S + 1/T$. The quality of the the ALOHA scheme is therefore simply $\frac{ST+1}{eST+1}$. We see that the ALOHA scheme then has a quality of 1 when the traffic is very bursty ($ST \ll 1$), i.e., it needs no more capacity than the M/M/1 scheme, and a quality $1/e$ when the traffic is very steady ($ST \gg 1$).

ALOHA systems with large populations have stability and control problems [3,8,9], but in the spirit of maintaining the simplest possible approximation we shall not deal with them.

In the *centralized* system described above, all messages have one common destination, even though their sources are distributed. When the traffic to be carried is between many terminal pairs we have a different problem, which we shall call the *network* problem. That is, in a network, both the sources of messages and their destinations are distributed.

In describing the centralized system we have implicitly assumed that all terminals can transmit with enough range to reach the station (i.e., we are not power limited), and that transmitting directly to the station is the best policy.

If the transmission range is not enough to span the distance from source to destination, the message will have to be received by some intermediate node and relayed towards its destination. That is, a message may need more than one *hop* in order to reach its destination. The intermediate node is often called *repeater*.

We have assumed that the centralized system is a one-hop system, but we shall explicitly treat the question of transmission range in networks, since it introduces an important tradeoff: a short transmission range makes more hops necessary, but reduces the interfering traffic. We shall see that choosing an appropriate range, as a function of traffic characteristics, will lead to the self-adjusting capability referred to in our title.

In section 2 we analyze networks assuming that the range of every transmission can be perfectly adjusted. In section 3 we analyze networks assuming the range of all transmissions must be equal. In section 4 we introduce the idea of power groups and show how it improves ALOHA systems. In section 5 we analyze hierarchical organizations of ALOHA systems with many levels.

2. Adjusting the Transmission Range

We assume that the transmission policy of all terminals is chosen to optimize the overall network performance. In order to analyze the tradeoff between range and interference we need a detailed model. We shall assume that our network covers a region of space that is large enough to make edge effects negligible. We shall also assume that terminals are placed everywhere with the same density, and that the terminal density is very high, so we may make all calculations as if we had a continuum of terminals. Other assumptions we adopt are [7]:

- (1) The rate of traffic exchanged between any two small geographic areas depends only on the size of the areas and the distance between them. The rate does not depend on the identity (i.e., location) of the areas or the direction from one to the other. That is, our network is homogeneous and isotropic in its statistical properties.
- (2) The access scheme used is slotted ALOHA. That is, we ignore the fact that the synchronization necessary for slotted ALOHA is hard to achieve in a network with long range transmissions and partially overlapping ranges.
- (3) The terminal's antenna is simple, and the signal propagates equally in all directions.
- (4) A transmission will not be bothered by other transmissions that are not within range of its (possibly intermediate) destination, but will be destroyed by any simultaneous transmission that takes place within range of its destination. A transmission will be successful whenever it is the only one within range of its destination. That is, we assume a definite range, beyond which no interference is felt. This is, of course, an abstraction of the real world, in which both successful reception and destructive interference are probabilistic events.

Consider a given terminal with a rate of s messages per slot destined to another terminal. A transmission will be successful only if there is no other transmission with enough range to interfere with it. Our terminal will have, therefore, to offer a total traffic of g messages per slot in order to succeed at a rate s , where g includes retransmissions of previously unsuccessful messages. Let G be the total offered traffic per slot heard at the destination. Assume that G is created by an infinite population

of terminals, and that the amount contributed to it by every source-destination pair is a Bernoulli process independent of the traffic offered by any other source-destination pair. Returning to our given terminal, whose contribution to G is minute, we must have $s = ge^{-G}$, where e^{-G} is simply the probability that no other message is transmitted in the slots used by our terminal. Summing over all transmissions heard at our destination we get

$$S_c = Ge^{-G} \quad (2)$$

where S_c denotes the total traffic successful at its destination and heard at our destination. This total traffic consists of messages with many different destinations, and the success of each message depends on what happens at its destination. But all these messages contend with *our* transmission for the use of the channel around *our* destination.

Equation (2) looks exactly like the equation describing a centralized slotted ALOHA system [6]. G and S_c do not, of course, depend on the transmission in question, and we can therefore say that any transmission sees an ALOHA system at its destination with a throughput equal to S_c , where the subscript on S_c stands for *contending*. If we *unnormalize* S_c and measure it in messages per unit time, we may use (1) and write the average delay per hop suffered by any message as follows:

$$T = \frac{1}{C - eS_c} \quad (3)$$

In the centralized case, interference always destroys both messages involved. In the network case analyzed here this is not necessarily true. Since the ranges of the transmission involved and their destinations may be very different, a collision of two messages at the first's destination will destroy the first, but may not bother the second at its destination. We shall use (3) for the delay in ALOHA networks, even though what happens at each destination is not equivalent to a closed, centralized ALOHA system; this is supported by [10] where the optimal transmission policy for ALOHA networks, given the hearing matrix, is shown to be identical to the optimal policy in centralized ALOHA systems. However, our goal here is to choose the optimum hearing matrix by choosing the transmission range.

The discussion so far applies to any network which is homogeneous and isotropic in a statistical sense. To be more specific let us assume that the terminals are distributed in an *infinite* two-dimensional region. That is, in a region whose size is much larger than the typical distance travelled by messages, so that edge effects can be neglected. Let S be the total traffic coming out of a unit area, and let $f(r)$ be the traffic density. That is, the traffic going from one small (source) area dA_s to another small (destination) area dA_d is given by $f(r)dA_s dA_d$, where r is the distance between the two small areas. We obviously have $S = \int_{r=0}^{\infty} f(r)2\pi r dr$ and $f(r)2\pi r/S$ is therefore the probability density function for the distance travelled by a message. N , the average distance travelled by messages, is given by $N = \int_{r=0}^{\infty} r f(r)2\pi r dr$. To calculate S_c , the total traffic contending at any destination, consider a message that must travel a distance of between r and $r+dr$. It will be heard at a given destination if it starts anywhere within a circle with radius r around that destination. We can then write

$$S_c = \int_{r=0}^{\infty} \pi r^2 f(r)2\pi r dr = \pi S \bar{N}^2 \quad (4)$$

Where $\bar{N}^2$ is the second moment of the distance travelled.

Assume now that the transmission range is chosen in such a way that every message will have exactly enough range to reach its destination in one hop. Substituting (4) in (3) we see that an ALOHA network in which every message reaches its destination exactly in one hop has the same delay-capacity relationship as a *centralized* ALOHA system carrying a total traffic $\pi S \bar{N}^2$.

The simplicity of (4) is a result of the assumption that power can be adjusted *exactly* to reach the destination. Two objections can be raised to this assumption:

- (1) Will a terminal always have enough power to reach its destination in one transmission?
- (2) Will the terminal have the capability to exactly adjust its power, and will it know the distance to its destination, on which this adjustment should be based?

These two objections are especially important in the environment of many cheap mobile terminals, which is exactly the environment which makes the ALOHA idea attractive.

We shall treat these objections later, but let us now ask another question: even if we can adjust the range so as to exactly reach the destination in one hop, is this a good policy? In [11] the question was posed thus: should we take giant steps, assuming we can? It was shown there that if, for a given C and traffic requirement, the delay per hop grows without bound as a function of the step size R , then there is an optimal step size, and steps should not be giant. We wish to find the optimal range policy as a function of traffic requirements, and for this we need the following:

Theorem 1: If a message must travel a distance X in k hops it should, in order to make the best use of the communication resources, do so in k equal hops, each of length X/k .

Proof: Whether we want to minimize T when S and C are given, or to minimize the necessary C when S and T are given, we must, in order to get the best system, minimize the total contending traffic at each destination. But this is equivalent to minimizing the total area at which a given message is heard. Let X_i be the size of the i -th hop, where $\sum X_i = X$. The area in which our message is heard is proportional to the $\sum X_i^2$. Minimizing the area at which our message is heard is therefore the following simple problem of constrained minimization:

$$\begin{aligned} &\text{Minimize } \sum X_i^2 \\ &\text{subject to } \sum X_i = X \end{aligned}$$

The solution of this minimization problem gives the equal step result stated in the theorem. □

Let us now consider the following family of policies which use a perfectly adjustable but limited transmission range. Given the maximum range R , the path of every message will be divided into the minimum number of equal hops. Which R will give the best overall system performance? Should we try to make R as large as possible? To answer these questions we must determine how S_c depends on R . Writing S_c as a function of R and the distribution of the distances travelled is a straightforward but cumbersome operation. However, the following bounds are simple to obtain:

Since $S_c(R)$ is a monotonic increasing function of R , an obvious bound is $S_c(R) \leq S_c(\infty) = \pi S N^2$. When R is very large all messages will reach their destination in one hop, so the equality here follows from (4).

Another bound, especially useful when R is small, can be obtained as follows: The total area covered by the several transmissions of a message that has to travel a distance r can be bounded from above by $\frac{r}{R} \pi R^2$ and $S_c(R)$ can therefore be bounded by

$$S_c(R) \leq \int_{r=0}^{\infty} \frac{r}{R} \pi R^2 f(r) 2\pi r dr = \pi R N S \quad (5)$$

Fig. 1 shows the two bounds and a hypothetical $S_c(R)$.

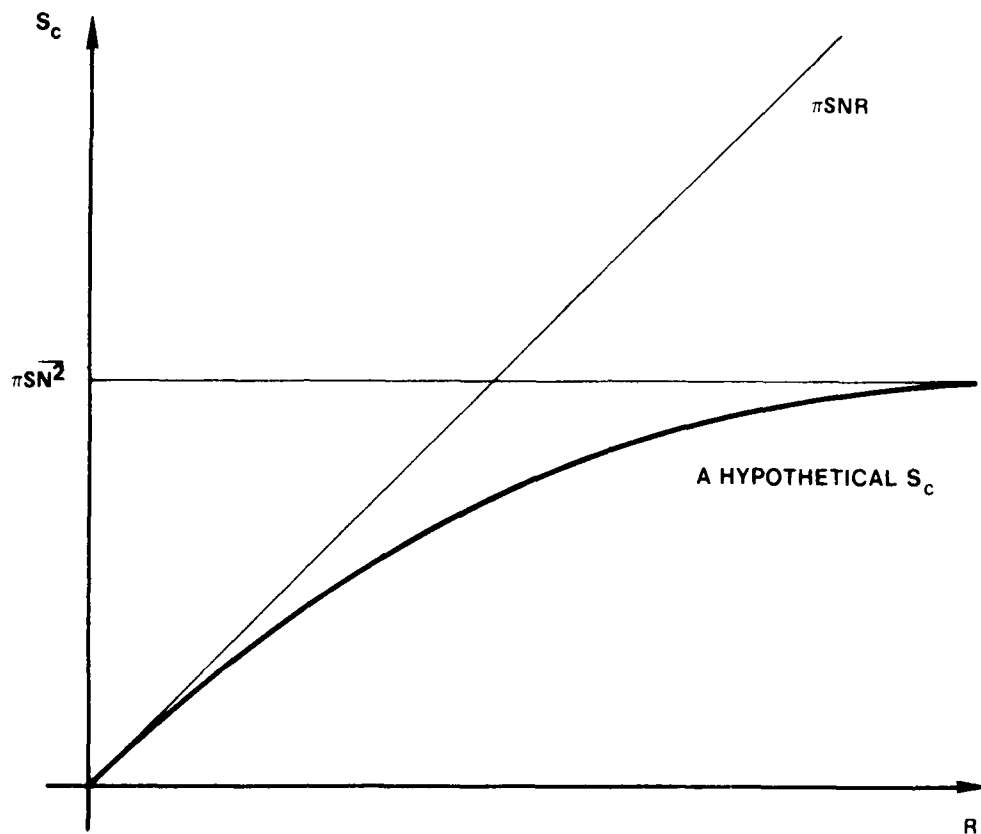


Figure 1. Two Bounds on S_c , the Total Successful Traffic Contending at Each Point.

We shall assume that the traffic to be carried is specified, that an acceptable delay is specified and that the goal of a good design is to make the necessary bandwidth as small as possible. The specification can be summarized by the dimensionless quantity N^2ST . When $N^2ST \ll 1$ we call the network, and the traffic, *bursty*, and when $N^2ST \gg 1$ we call the network *steady*.

For small R we can use the bound of (5) as an approximation for $S_c(R)$, and we will combine it with N/R as an approximation for the average number of hops per message, to get the following approximate expression for the delay

$$T = \frac{N/R}{C - e\pi SNR}$$

Inverting we get

$$C = e\pi SNR + \frac{1}{T} \frac{N}{R} \quad (6)$$

and from this approximate expression for C we can get that the optimal R (i.e., the R that minimizes

communicate with the station using ALOHA. All communications will use the full capacity of the channel. Repeaters may be necessary in order to extend the range of transmission, but we shall assume this is not a problem, and shall only be interested in introducing repeaters in order to improve system performance. That is, to lessen the delay when S and C are given, or lessen the capacity necessary when S and T are given.

Should all groups of terminals be of the same size? To minimize contention in the bottom level all ALOHA subgroups should carry the same traffic, i.e., a symmetric bottom level is best. But in order to reduce the contention in the top level we should have as much asymmetry as possible. The best top level will consist of one repeater forwarding all the traffic to the station without any conflict. But such a two-level system will not help us, because its bottom level itself is equivalent to the one-level system we set out to improve.

Since two-level systems are introduced in order to reduce the contention in the bottom level we shall assume that whenever two levels are better than one, the traffic is evenly divided between groups. Gitman [14] introduced such a scheme and calculated the capacity of two-level systems. He assumed all terminal groups can use the same channel without interference, that a terminal cannot be successful when its repeater is talking to the station, i.e., the repeater cannot talk and listen at the same time, and that a terminal may be influenced by other repeaters talking. The largest capacity is obtained when the terminal is influenced only by its own repeater, and when there are only two repeaters. But even then the capacity obtained is less than $1/2$. The reason for this is the following: Let S be the total throughput in the system. Let G be the total offered traffic in the top level, that consists of two repeaters. G is larger than S because it includes the retransmissions of messages that have been previously transmitted unsuccessfully. In a system with capacity C the slot size will be $1/C$ since we have chosen an information unit such that the packet length is one. The throughput and the offered traffic per slot will then be S/C and G/C . We shall assume that the traffic per slot offered from each of our two repeaters is a Bernoulli process, i.e., a discrete Poisson process, which is independent (!) of the traffic offered by the other repeater. A transmission from a repeater will be successful only if the other is not transmitting in the same time slot. Calculating the total success rate in the top level we get

$$\frac{S}{C} = 2 \frac{G}{2C} \left(1 - \frac{G}{2C} \right) \quad (28)$$

where $G/2C$ is the total traffic offered by each one of the two repeaters, and $1 - G/2C$ is the probability that a packet succeeds. In order to achieve $S/C \approx 1/2$ we must have $G/C \approx 1$, so that each of our two repeaters is talking half the time. It is impossible to feed such a talkative repeater from an infinite population of terminals, because the maximum success rate of each of the two groups is $(1 - G/2C)/e = .184 < .25$.

The maximum throughput of such a two-level system is given by the following set of equations

$$S = G \left(1 - \frac{G}{2C} \right)$$

$$\frac{S}{C} = 2 \left(1 - \frac{G}{2C} \right) \frac{1}{e}$$

from which we get that the maximum S/C is equal to 0.465. So even though we cannot achieve the full capacity of a two terminal system we do get an improvement over a one-level ALOHA

improvement even though the considerate strong group carries more of the traffic, moves with smaller steps and uses less of the channel relative to the weak group.

For a summary of the optimal range and the necessary capacity in various two-dimensional networks see Table I.

Table I

Best Transmission Range and Needed Capacity for Two-Dimensional Networks

Organization		Range	Capacity
M/M/1		R_0	C_0
ALOHA (one group)		$.607R_0$	$1.647C_0$
ALOHA (two groups, same range)		$.729R_0$	$1.372C_0$
ALOHA (two groups, separate ranges)	selfish	$.774R_0$	$1.292C_0$
	considerate	$.782R_0$	$1.279C_0$

$$R_0 = \frac{1}{\sqrt{\pi ST}} \quad C_0 = 2\sqrt{\pi N^2 S/T}$$

5. Multi-Level ALOHA

Until now we have always assumed our ALOHA systems have an *infinite* population. Let us now consider a slotted ALOHA system with a finite number, m , of equally talkative terminals. Assuming the traffic offered by any terminal at a given slot is independent of the traffic offered by other terminals or at other time slots we can simply see [2] that such an m -terminal ALOHA system can successfully utilize a fraction of the time slots equal to

$$\left(1 - \frac{1}{m}\right)^{m-1} \quad (26)$$

The rest of the time slots will be wasted on destructive interference, or will be left unused even with some messages waiting for transmission. This last occurrence is necessary in an optimized system to ensure the fraction of slots wasted on collisions is not too large.

When m is very large, Eq. (26) states that the maximum utilization of an infinite population slotted ALOHA is $1/e$, which is the expression we used before. But when m is finite the ALOHA system can do better. The best case is when $m=2$, and the maximum possible utilization is then $1/2$. One could also talk about an ALOHA system with only one terminal, that can use all time slots without any wasteful collisions, but this case is of no interest.

In analogy to (1) we shall model the delay of a finite population slotted ALOHA system by

$$T = \frac{1}{C - S/U_m} \quad (27)$$

where U_m is the maximum possible utilization of an m -terminal system, as given by (26)

Since ALOHA systems with a small population have better utilization and smaller delay than ALOHA systems with a large population, one is led to the following hierarchical scheme. Divide the very large terminal population into a small number of groups. Assign a repeater to each terminal group. Each group will communicate with its repeater using ALOHA, and the repeaters will

Can this result be improved if messages from the two groups will not travel the same average distance? Let N_1 and N_2 be the average distance travelled by messages from the strong and weak group, respectively. Substituting N_1 for N in (23) and N_2 for N in (24) we have T_1 and T_2 . Our goal now is to minimize T subject to $N_1S_1 + N_2S_2 = NS$ and subject to $S_1 + S_2 = S$. It is easy to see that T is minimized when $\frac{N_1S_1}{N_2S_2} = \frac{1}{b^2} = 1.261$, and that the minimal T is given, once again, by (25). That is, the added flexibility of giving each group of messages a different average distance does not lead to a better network!

It is interesting to note that $\frac{T_1}{T_2} = \frac{N_1}{N_2}$ but that $\frac{S_1T_1}{S_2T_2} = \frac{1}{b^2} = 1.261$. That is, we can choose the ratio between T_1 and T_2 at will (by adjusting N_1/N_2) but the contribution of the strong and weak group to the average delay and to the average number of messages in the network will always, in an optimized system, be in a fixed ratio.

Let R_1 and R_2 be the maximum hop size in the strong and weak group. Using (7) we see that when T is minimized $\frac{R_1}{R_2} = b = .793$. That is - the strong group carries much more of the traffic, and even though it has more bandwidth available, it uses smaller hops.

When choosing S_1, S_2, N_1 and N_2 in order to optimize the two-dimensional network with two groups, we have assumed that the strong group is selfish. But we saw before that a better overall system can be obtained if the strong group is not absolutely selfish, and does not use the channel to its utmost. How considerate should the strong group be in a network?

The average delay in the strong group can be written (when $N_1 \gg R_1$) as

$$T_1 = \frac{N_1/R_1}{C - e\pi R_1 N_1 S}$$

Here we cannot use (8) because when the strong group is considerate it will use a smaller range R_1 than the range used by a selfish group.

The weaker group does not bother anyone, and should use what is available to it to the utmost. Let b denote, once again, the fraction of capacity available to the weak group. (b is now a design variable, parametrizing the amount of consideration shown by the strong group). To the weak group we can apply (8), and we then get $T_2 = 4e\pi \frac{N_2^2 S_2}{b^2 C^2}$. Our goal is to minimize $ST = S_1T_1 + S_2T_2$ by choosing S_1, N_1, R_1, S_2, N_2 and R_2 subject to $S_1 + S_2 = S$, $S_1N_1 + S_2N_2 = SN$. When choosing S_1, N_1 and R_1 we also determine b . To see this let us denote by G the total traffic (per time slot) offered by the strong group which is heard at any given point. G can be determined by equating the following two expressions for the success rate of strong messages at any local ALOHA system. $Ge^{-G} = \pi R_1 N_1 S_1 / C$. b , the fraction of time slots left free by the strong group, is given by $b = e^{-G}$. T obviously depends on S_1, N_1, S_2 and N_2 only via the products S_1N_1 and S_2N_2 . The results of choosing the best S_1, N_1 and S_2, N_2 , for a given R_1 , can be most simply written in terms of G

$$T = 4e\pi \frac{N^2 S}{C^2} \frac{1}{4eGe^{-G} \left(1 - eGe^{-G} \right) + e^{-2G}}$$

The G which minimizes T can be found by numerically solving the equation $dT/dG = 0$, and is given by $G = 1.79$. b is then equal to $e^{-G} = .836$ and the quality of this best two-group two-dimensional network is then .782. In this network, with a considerate strong group, we have $N_1S_1/N_2S_2 = 1.380$ and $R_1/R_2 = 0.704$. Comparing with the selfish case we see that consideration leads to an overall

strong group is $e-1$ times the traffic contributed by the weak group.

Until now we have applied the idea to partially coordinated groups (i.e., power groups) to centralized ALOHA systems. How can it be applied to networks? In our analysis of ALOHA networks we have used the transmission power to control range. We shall now assume that the division into groups is done by means which are independent of power so that transmission range can still be freely chosen. We shall also assume that the policy of assigning transmission power is independent of position, and that the density of both strong and weak sources is high and uniform.

One simple way to improve ALOHA networks by using groups is the following: The same transmission range will be chosen for both *strong* and *weak* transmissions, and the partial coordination between them will simply improve the local ALOHA system. From (18) we get that the maximum local utilization of a two-group ALOHA system is 0.531. Substituting this in (8) we see that by using two groups with the same range the quality of a two-dimensional network can be improved from $\sqrt{.367}=.607$ to $\sqrt{.531}=.729$. In one-dimensional networks the quality is equal to the local utilization and using two groups will improve both from .367 to .531.

We see that since networks of high dimensionality are less sensitive to the limited utilization of the ALOHA scheme it is harder to improve them by introducing a better scheme.

The capability to divide terminals into two partially coordinated groups can lead to a greater improvement of ALOHA networks (in two or more dimensions) if transmission range is chosen independently for the two groups.

Let us consider a two-dimensional network and assume at first that the average distance travelled by transmissions from both strong and weak groups is equal to N . We shall also assume that if a message needs more than one hop then all of its hops will be strong or all of them will be weak. Let S_1 and S_2 be the traffic density of the strong and weak group, and let T_1 and T_2 be the average delay suffered by messages from the strong and weak group respectively. In a heavily loaded system, if the strong group is absolutely selfish it will utilize the full channel in the way best for it and we then get from (8) that T_1 and S_1 satisfy

$$T_1 = 4e\pi \frac{N^2 S_1}{C^2} \quad (23)$$

The local utilization of the strong group, when optimized for heavy traffic, is $1/2e$. It is easy to calculate that the strong group leaves then a fraction $b=.793$ of the time slots unused, and these slots are available for the weak group. That is, the capacity available to the weak group is bC . Using (8) we get that

$$T_2 = 4e\pi \frac{N^2 S_2}{b^2 C^2} \quad (24)$$

T , the message delay averaged over all messages, from both groups, is given by $TS = T_1 S_1 + T_2 S_2$, and our goal is to minimize T by choosing S_1 and S_2 subject to $S_1 + S_2 = S$. It is simple to see that T is minimized when $S_1/S_2 = 1/b^2 = 1.261$ and is then given by

$$T = 4\pi \frac{e}{1+b^2} \frac{N^2}{C^2} S \quad (25)$$

The quality of this two-group network is therefore $\sqrt{(1+b^2)/e} = .774$.

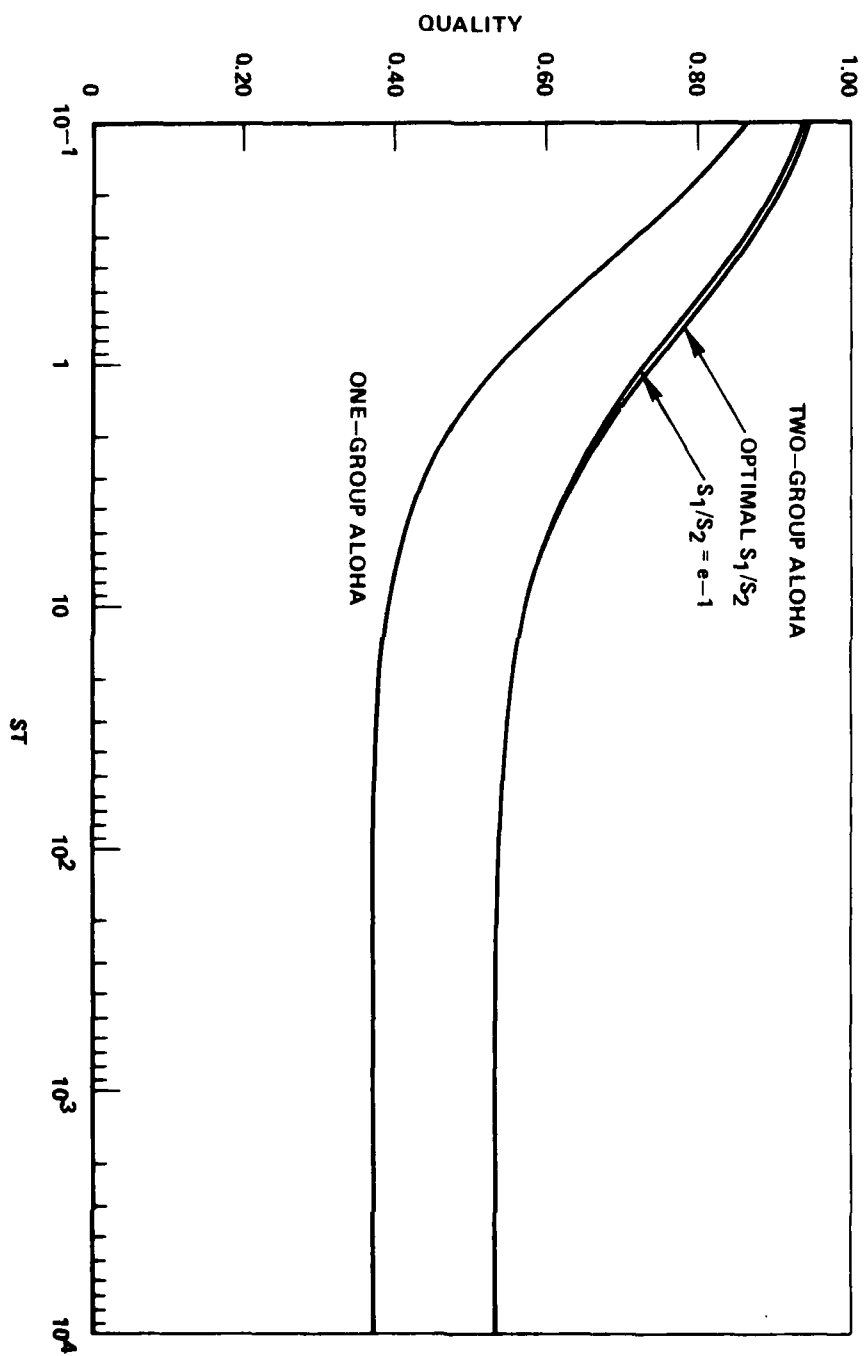


Figure 4. Quality of ALOHA With Two Power Groups.

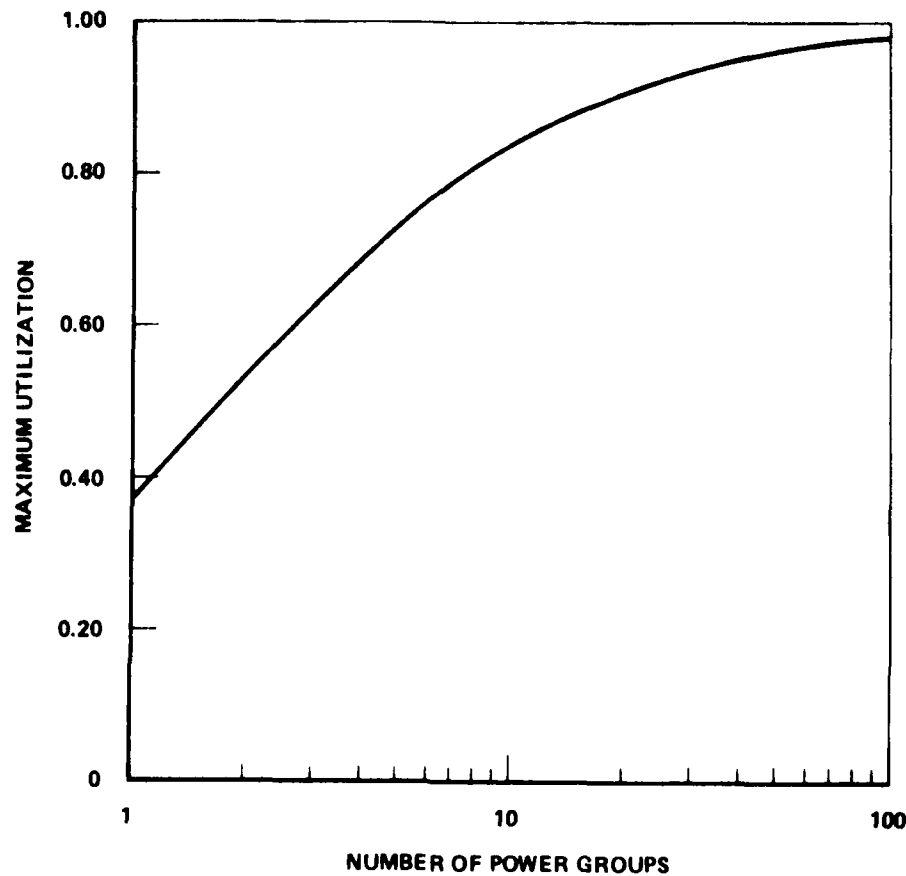


Figure 3. Maximum Utilization of ALOHA With Power Groups.

This expression for T is a weighted sum of two delay terms. The first term corresponds to the delay in the strong group, (which can ignore the weak group and that behaves like an ALOHA system with capacity C and traffic S_1 .) The communication capacity available to the weak group is $Ce^{-S_1/C}$ since this is the portion of the channel left unutilized by the strong group. The second term corresponds to the delay in an ALOHA system with this reduced capacity carrying a traffic S_2 .

With a given C and with a given total traffic $S=S_1+S_2$, which S_1 and S_2 will give the minimum delay? The best S_1 and S_2 as a function of load can be found numerically: Fig. 4 gives the quality of the two-group ALOHA system thus optimized. When the system is only lightly loaded, S_1/S_2 is only slightly larger than one. When the load grows this ratio also grows, and when the system is driven to its maximum utilization S_1/S_2 goes to $e-1$, as given by (19). Also shown in Fig. 4 is the quality of a two-groups system in which the ratio between S_1 and S_2 was always chosen by (19), which is the optimal choice at heavy traffic. We see that the improvement gained by optimizing the ratio between S_1 and S_2 as a function of load is negligible, and that a very good two-group centralized ALOHA system can be obtained by simply splitting the terminal population so that the traffic contributed by the

$$e^{-(1+e)} = 0.531 \quad (18)$$

and that this utilization is achieved when

$$\frac{S_1}{S_2} = e-1 \quad (19)$$

The above treatment can be generalized to many groups. Assume that the terminals are divided into r progressively weaker groups where the following is true: A message will never be bothered by transmissions from weaker groups, and will always be destroyed by transmission from its own group or from a stronger group. We then have:

Theorem 5: Let V_r be the maximum utilization of a slotted ALOHA system whose infinite population is optimally divided into r groups, with the above assumption on immunity to some cases of interference. Then V_r satisfy the following recursion relation:

$$V_{r+1} = e^{-(1+V_r)} \quad (20)$$

Proof: In analogy to the two group case we can write

$$S_1 = G_1 e^{-G_1} \quad (21)$$

$$S_2 = G_2 e^{-G_2} e^{-G_1}$$

$$S_r = G_r e^{-G_r} \dots e^{-G_1}$$

V_r is obtained when $S = S_1 + S_2 + \dots + S_r$ is maximized by varying the G_i . Since no transmissions from weaker groups will ever influence the strongest group we can optimize their throughput separately, and (21) will then reduce to

$$S_1 = G_1 e^{-G_1} \quad (22)$$

$$S_2 + \dots + S_r = V_{r-1} e^{-G_1}$$

The optimal G_1 is then easily found to satisfy $G_1 + 1 = V_{r-1}$, and substituting this G_1 into (22) we get (20).

The sequence V_r , whose first portion is shown in Fig. 3, is a monotonic increasing sequence converging (slowly!) to 1. This is not surprising since when we have a large number of groups most collisions will be between messages from different groups, and one of the messages will be successful.

Having a large number of groups with the clear separation assumed in Theorem 5 may be impractical. But having two groups is reasonable, and we shall discuss this case in some detail.

Eq. (18) gives the maximum utilization of a two-group ALOHA system. What will be the delay in this system? Returning now to our custom of measuring S in messages per unit time (and not per slot), we shall model the delay by

$$T = \frac{S_1/S}{C - eS_1} + \frac{S_2/S}{C e^{-G_1} - eS_2}$$

where G_1 is given by $S_1 = G_1 e^{-G_1}$

4. Capture, Power Groups and Partial Coordination

In the models of ALOHA systems presented so far we assumed that in the case of interference, both messages will be destroyed. But if the colliding messages vary greatly in received power, the receiver may be able to receive the stronger one correctly even in the presence of the other, weaker, signal. The receiver is then said to *capture* the stronger signal. The capability to capture some messages will obviously improve every ALOHA system. Let us first see the resulting improvement in a centralized ALOHA system, where all messages have one common destination. Roberts [6] proposed and analyzed a capture model in which the power differences resulted from different distances to the common destinations. Our approach is different. We shall assume that the terminal population is split into two groups, that one group is transmitting with more power than the other, and that this splitting is purposely done in order to improve system performance. In order to abstract the geometric details out of the model, we shall adopt the following assumption [13]: The power of the two groups is significantly different. When two transmissions from the same group occur simultaneously, they will always destroy each other. When one strong transmission and any number of weak transmissions compete for the ear of the common station, the strong one will always be captured successfully. This separation into groups introduces, therefore, a partial coordination into the random world of ALOHA.

It may be possible to achieve such a coordination between groups by techniques that do not rely on a power difference between them. A distinctive preamble, for example, may allow a terminal to successfully receive a transmission from one group, which we shall call *strong*, even in the presence of transmissions from the weak group. In a system which is not perfectly slotted, the first of two interfering signals of equal strength to arrive at a receiver may survive the collision and be successfully received. From now on strong and weak should not therefore be taken literally - they do not necessarily refer to transmission power, but simply characterize the group of transmissions likely to win or lose when competing with the other group.

What will be the resulting improvement if we introduce groups into a heavily loaded ALOHA centralized system? If the strong group is *selfish* it can ignore the weak group, and use the channel as much as possible. The strong group will then successfully utilize $1/e = .367$ of the slots, and will leave .367 of the slots free. (In addition, .276 of the slots will be wasted on collisions). The weak group can utilize at most $1/e$ of what is left free for it, i.e., it can utilize $1/e^2 = 0.135$ of the slots, and the total rate of success by both groups will be 0.503.

The channel can be better utilized if the strong group will not be so selfish. To see this let us now consider the division into groups as a design parameter.

Assume that we have an *infinite* population of terminals, and that each terminal contributes only a minute fraction of the total traffic. While we have spoken of strong and weak terminals, the important design question is *not* the identity of terminals in each group but the portion of the traffic in each group. If we have an extremely heavy load our goal is to find the division into groups that will allow our system to utilize the greatest portion of the communication resource available. Let G_1 and S_1 be the total offered traffic and the rate of success of the strong group, G_2 and S_2 the corresponding values for the weak group. For simplicity we shall assume in this section that S and G are measured per slot size. Using our standard assumption, that the total traffic offered by a terminal is a Bernoulli process, independent of the traffic offered by all other terminals, we can write

$$S_1 = G_1 e^{-G_1}$$

$$S_2 = G_2 e^{-G_2} e^{-G_1}$$

Choosing G_1 and G_2 in order to maximize $S = S_1 + S_2$ we find that the utilization of a system with two groups is

From (17) we get the following:

Theorem 4: When the traffic is very steady the transmission range that is optimal when all transmissions must have the same predetermined range is equal to the optimal maximum range when range can be perfectly adjusted.

Proof: When the traffic is very steady and R is small, the bound of (17) is a good approximation for $S_c(R)$. Using this expression for S_c we can continue as in the proof of Theorem 2.

□

This theorem is very intuitive: When the optimal step size is small, the capability to adjust transmission range is not important, since the overshoot will be small.

It immediately follows that the network quality and the local utilization that were used in Theorem 2 to characterize the optimal network for very steady traffic when the range is perfectly adjustable will also characterize the optimal network when the range must be predetermined.

When the range is perfectly adjustable the one-dimensional network was a special case, in which giant stepping was appropriate. When the range must be predetermined we see from (16) that S_c increases with R . When all transmissions have a range equal to R it must, therefore, be limited even in the one-dimensional network.

When the traffic is very bursty ($N^*ST \ll 1$) we expect R to be larger with respect to N , and shall then use the bound given in (16) as an approximation for S_c . When R is large we also assume no message takes more than two hops and we approximate H , the average number of hops taken by a message, by

$$H = 1 + \text{Probability (distance travelled} > R)$$

The capacity necessary can then be approximated by

$$C = evR^*S + \frac{H}{T}$$

and the R that will minimize C is now given by solving the following equation:

$$nevSTR^{n-1} = \text{Probability density (distance travelled} = R)$$

For a very large R it is reasonable to assume that the probability density of the distance travelled is monotonic decreasing and this equation will then have a unique solution. If, for example, the distribution of distances travelled is exponential we get the following approximate equation defining the optimal R in a bursty system: $R/N = \ln(1/veNSTR^{n-1})$.

When considering centralized systems we can say that the ALOHA scheme is good when the traffic is bursty and bad when the traffic is steady. This statement is true in general for ALOHA networks too. But networks have self adjusting property - by controlling the maximum transmission range and reducing it when the traffic is steady we can make ALOHA networks (in more than one dimension) suffer less from destructive interference than the ALOHA centralized system.

In the next two sections we shall consider two other ideas that can improve a centralized ALOHA system and see what they can contribute to ALOHA networks.

In Theorem 2 we assumed $n > 1$. The reason for this is that Theorem 1 can be generalized only for the case $n > 1$. When dealing with a one-dimensional ALOHA network we get

Theorem 3: S_c , the amount of contending traffic heard at a point, is equal to $2NS$, and is independent both of the need to break message paths into several hops and of the policy of implementing such a break, as long as the policy is applied everywhere in the same way. That is, as long as a message path of a given length will be broken in the same way, wherever it originates.

Proof: Consider a message that must travel a distance X , and let X_i be the length of its i -th hop, where $\sum X_i = X$. The i -th hop will be heard at a given point if the path of the message is so placed that the i -th hop starts within X_i of that given point, on either of its sides. Adding the contribution of all the hops we see that a message whose total path length is between X and $X+dX$ will always contribute to S_c an amount proportional to $2X$. In this one-dimensional network we have $S = \int_0^\infty f(x) dx$ and $N = \frac{1}{S} \int_0^\infty x f(x) dx$. S_c is therefore given by $S_c = \int_0^\infty 2x f(x) dx = 2NS$. □

In one-dimensional networks, if range can be perfectly adjusted we should, therefore, giant-step whenever possible. Even when the traffic is very steady there is no reason to limit the step size, since no decrease in S_c will follow. One-dimensional ALOHA networks have a local utilization and a network quality both of which are equal to $1/e$.

Theorem 2 answers the question of the optimal transmission range when the traffic is very steady. This is satisfying because ALOHA has an efficiency problem exactly when the traffic is steady. When the traffic is bursty there is little need for improving the ALOHA network. When range is perfectly adjusted the range limit R grows when the traffic becomes bursty, and when the traffic is very bursty giant stepping is the best (for all n). That is, each message should be transmitted with enough range to reach its destination directly (in one hop). These general conclusions change once we consider networks in which range cannot be perfectly adjusted, as we shall now do.

3. Using A Fixed Range

Assume that terminals cannot adjust the range of their transmissions, and that all transmissions, by all terminals, must have a fixed range R . Since the range of all transmissions is fixed and constant, some messages will overshoot their destinations. The amount of traffic contending at every point will therefore be larger now than it was when range was perfectly adjusted. S_c will depend on R in a way that involves the distribution of distances travelled by a message, but the following bounds are simple to obtain:

In an n -dimensional ALOHA network

$$vR^n S \leq S_c \quad (16)$$

because at every point we hear at least the first hop of all messages originating within R . In analogy to (11) we get

$$S_c \leq vSR^{n-1}(N+R) \quad (17)$$

because the average distance actually travelled by a message when the transmission range is predetermined at R is at most $R+N$.

Theorem 2: Consider an n -dimensional ALOHA network carrying very steady traffic, where $n > 1$. Assume that the transmission range can be perfectly adjusted, but only up to a maximum range R . If R can be optimized freely (i.e., made as small as necessary) then each transmission will see an ALOHA system whose local utilization is $1/ne$ and the network quality will be $1/e^{1/n}$.

Proof: The volume of an n -dimensional sphere with radius R is vR^n , where v is a constant depending only on n (when $n=2$ $v=\pi$). Theorem 1 will be valid for any $n > 1$. That is, if a message must travel more than R it should do so in the minimum number of equal hops. In analogy with (5) we therefore get

$$S_c(R) \leq vSNR^{n-1} \quad (11)$$

When the traffic is very steady and when $R \ll N$ this bound is a reasonable approximation for S_c and we get the following estimate for the capacity necessary when S and T are given:

$$C = evSNR^{n-1} + \frac{1}{T} \frac{N}{R} \quad (12)$$

The R that minimizes C is given by

$$R = \left[\frac{1}{evST(n-1)} \right]^{1/n} \quad (13)$$

and using this optimal R we find that the capacity necessary is

$$C = \frac{N}{T} n \left[evST(n-1)^{1-n} \right]^{1/n} \quad (14)$$

For that n -dimensional M/M/1 network we get a set of equations very similar to (12)-(14), but in which 1 is substituted for e . (Compare for example (8) and (10) in the two-dimensional case.) From (13) we see that the optimal R in an n -dimensional ALOHA network is smaller than the optimal R in an n -dimensional M/M/1 network by $1/e^{1/n}$. As long as this smaller R is consistent with our model we can derive from the dependence of C on e shown in (14) the quality part of the theorem.

The local utilization is, by definition, equal to

$$\frac{S_c}{C} = \frac{vSNR^{n-1}}{C} \quad (15)$$

and substituting (13) into (15) we find that when the optimal R is used the local utilization is $1/ne$.

Theorem 2 can be immediately generalized to the situation in which the antenna carried by terminals is somewhat directional. Assume the antenna radiates into a cone, which takes a fraction α of the sphere. This is, of course, a gross simplification of the real radiation pattern, but is consistent with our simple modeling of transmission range. If we compare the case of an omni-directional antenna to this case of an α -directional antenna we find that, with any transmission policy, the total interfering traffic at any point is smaller by a factor α . The optimal R for steady traffic, given by (13), will become larger by $1/\alpha^{1/n}$ (we shall not have to push so much towards small R), and the necessary capacity of (14) will become smaller by $\alpha^{1/n}$. But when we compare an α -directional ALOHA network to an α -directional M/M/1 network we find that the local utilization and the network quality in the optimized structure will remain as stated in Theorem 2. An improved technology (i.e., directionality) will help both the ALOHA network and the M/M/1 network. But whenever they use the same technology a comparison between them will show the inherent cost due to the random access aspect of the ALOHA network, and this inherent cost is $e^{1/n}$.

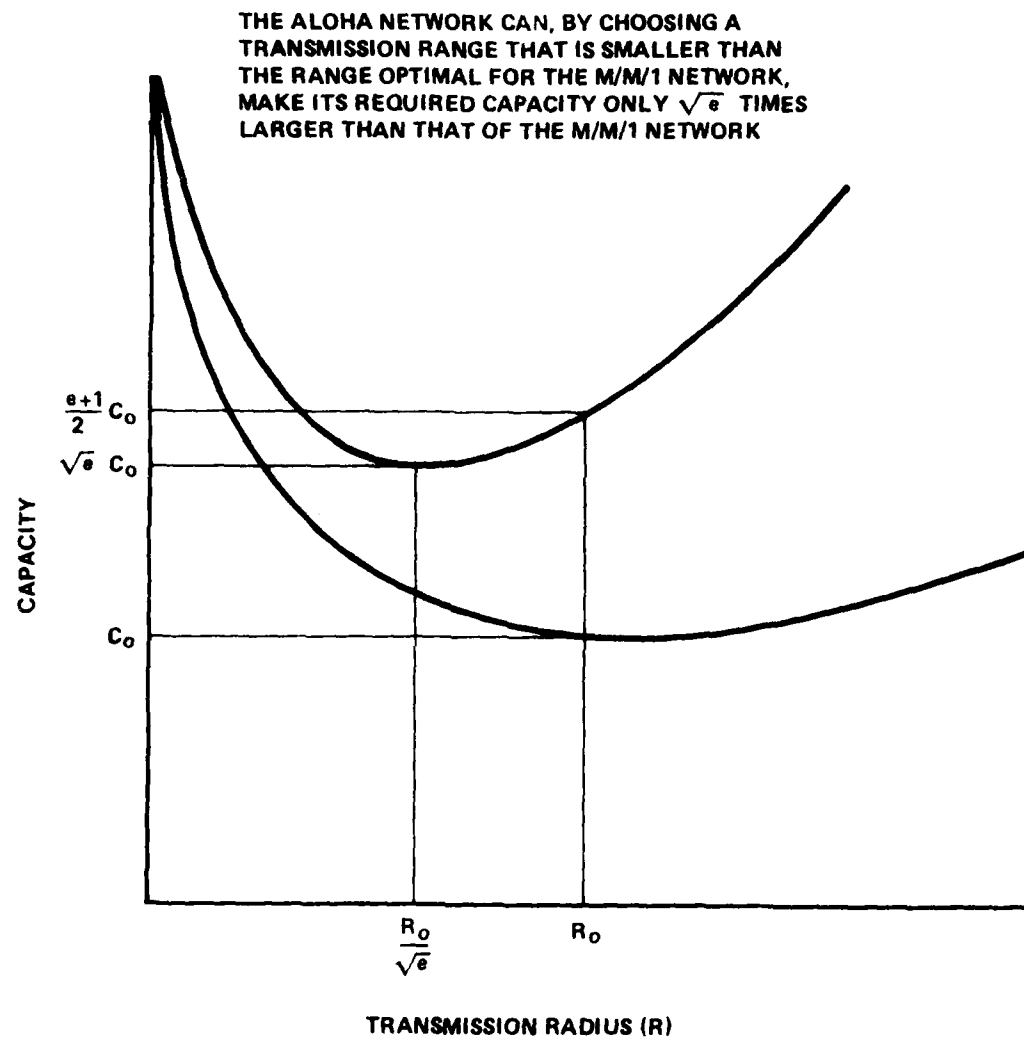


Figure 2. Capacity Necessary for Very Steady Two Dimensional Networks.

best possible M/M/1 network scheme is in general a function of S , T and the distribution of distances travelled. For very steady traffic we get, in analogy to (7), that the optimal R is given by

$$\frac{R}{N} = \frac{1}{\sqrt{\pi N^2 S T}} \quad (9)$$

and when using this R the capacity necessary is

$$CT = 2\sqrt{\pi N^2 S T} \quad (10)$$

Dividing (10) by (8) we get that the quality of heavily loaded two-dimensional ALOHA network with the optimal step size is $1/\sqrt{e} = .607$! How did we get this dramatic improvement over the heavily loaded centralized ALOHA system, whose quality is $1/e = .367$?

We may say that every message *sees* at its destination an ALOHA system whose utilization, which we shall call *local* utilization, is S_c/C . When the traffic is very steady and when the optimal R is used we get from (7) that every transmission *sees* an ALOHA system whose local utilization is $1/2e$, i.e., half the maximum possible utilization of an ALOHA system. The quality of a *centralized* ALOHA system with this local utilization is .68. It is only at much higher utilizations (closer to $1/e$) that the quality of a centralized ALOHA system goes down to $1/e$. The need for several hops will bring the quality of the ALOHA network down, from .68 to .607. We see therefore that by choosing the optimal R as a function of burstiness our ALOHA network has gained a *self-adjusting capability*, and it will not allow itself to be pushed to higher loads, where it is really bad.

From (8) and (10) we see that two-dimensional networks with the optimal R show an economy of scale when very steady: for a given T , the necessary C grows only like $\sqrt{S}$.

Comparing (7) and (9) we see that the optimal transmission radius R in a steady ALOHA network is smaller than the optimal R in an M/M/1 network by a factor $1/\sqrt{e}$. The optimal R in both networks goes to zero as the traffic becomes very steady. We have implicitly assumed that there always is a terminal at the end of the hop that can receive our message and forward it. But if R becomes too small there may not be a terminal so conveniently situated. If R becomes even smaller, our terminal may not be able to communicate with any other terminal, and the network may become disconnected. Kleinrock and Silvester [12] treat this issue explicitly, while calculating the optimum transmission range with a different objective: obtaining the maximum throughput from the given channel, assuming infinite delay is acceptable. We shall not treat this issue here, but our assertion about the self-adjusting capability of ALOHA networks must be qualified.

Consider once again an ALOHA network and an M/M/1 network, both carrying the same very steady traffic. If it is practical for the ALOHA network to choose the optimal R according to (7) then it will need only $\sqrt{e}$ times more capacity than the optimal M/M/1 network, i.e., its quality will be $1/\sqrt{e}$. But if R cannot be made so small, the quality of the ALOHA network will go down. If the ALOHA network is constrained to use the same R as the optimal M/M/1 network then its local utilization will be $1/(e+1) = .269$ and its quality will be $2/(e+1) = .538$. If both the ALOHA and the M/M/1 networks carry a very steady traffic but are constrained to use an R that is much larger than the one given by (9) then the local utilization of the ALOHA network and its quality will be $1/e$.

Fig. 2 sketches the dependence of the necessary capacity on the transmission range, in the two-dimensional ALOHA and M/M/1 networks.

Our treatment of two-dimensional networks can be summarized and generalized to n -dimensional networks as follows:

the necessary C for given N, S and T is given by

$$\frac{R}{N} = \frac{1}{\sqrt{\pi e N^2 S T}} \quad (7)$$

While we use the term optimal R , equation (7) actually determines the optimal value for the *maximum* transmission range. Given the distance a specific message must travel, R determines the necessary number of hops, and the transmission range of all hops is then chosen according to Theorem 1. The capacity necessary when using the optimal R can be obtained from (6) with the use of (7); it is given by the following relation between CT and N^2ST , both of which are dimensionless quantities,

$$CT = 2\sqrt{\pi e N^2 S T} \quad (8)$$

When the traffic is very steady (i.e., when $N^2ST \gg 1$) (7) says that the optimal R will be much smaller than N . The approximations made when writing (6) are consistent with this result, which is also quite intuitive: Consider a steady system with a given S and a large T . When we are willing to tolerate a large T the number of hops can be large, and we can therefore choose a small R . Each message will then be heard only in a narrow strip along its path, so S_c will be small, and the necessary bandwidth will therefore also be small. When the traffic is very bursty we get from (7) that R is much larger than N . This is again very intuitive - when the traffic is bursty there is little contention and therefore almost nothing is gained by forcing a message to undergo more than one hop. But the exact value given by (7) is not meaningful when the traffic is bursty, because the approximations used when writing (6) are not valid when R is large.

A general conclusion that emerges is that in a two-dimensional network it is better to limit the transmission range even if our terminals can adjust their range exactly and have no power limitation. This *voluntary* limiting is especially important when the traffic is very steady, and the optimal range limit R is then given by (7).

How shall we define the *quality* of networks? Clearly one should *not* compare a network to one huge centralized M/M/1 system that carries all messages to one common destination because practical networks have an advantage over centralized systems: The same capacity can be used in different regions of the network to successfully transmit different messages at the same time. That is, network capacity can be spatially reused.

A common measure used to characterize access schemes is the maximum utilization they can make of the given communication resources. This maximum utilization is sometimes called capacity, especially by authors whose variables are normalized by the slot size, and who therefore do not explicitly mention the channel bandwidth. We use the word *capacity* to describe an amount of communication resources (i.e., the number of bits or messages that can be transmitted per second) and *utilization* to denote the useful fraction of that capacity.

The quality of a very steady *centralized* system, as defined by us [7], is equal to its maximum utilization. But utilization is not a good measure for networks with a continuum of terminals since utilization can be arbitrarily increased by spatial reuse, i.e., by limiting the transmission range.

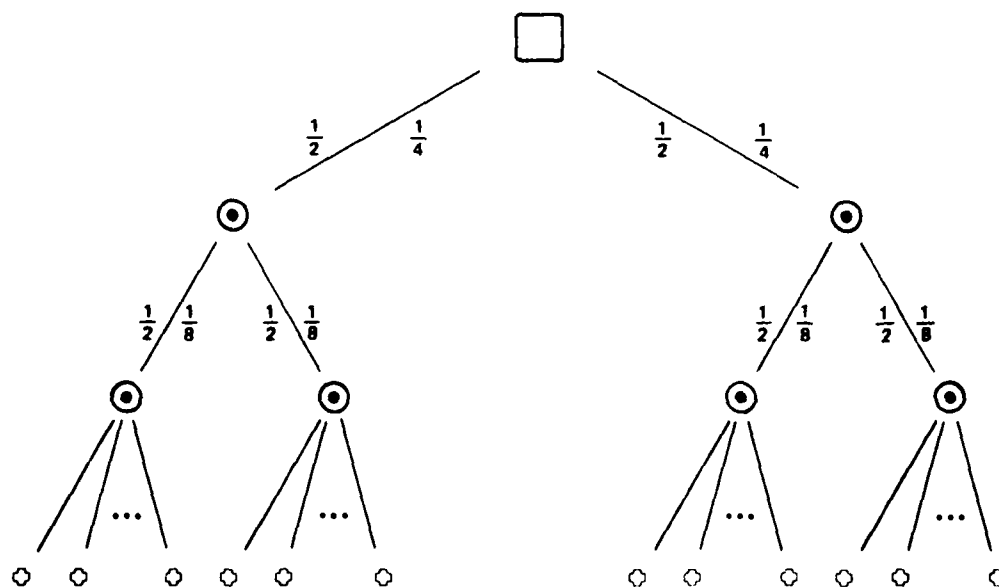
It seems that every network organization must address the question of how to coordinate every transmission with at least all the traffic that is heard at its destination. Since the best possible system will coordinate this traffic perfectly, we shall compare all networks to the network that uses the same technology (i.e., omni-directional antennas) but that somehow achieves perfect coordination between the traffic contending at every point, and in which transmission ranges are chosen optimally. We shall define the quality Q of any network to be the inverse ratio between the capacity necessary for it when S and T are given and the capacity necessary in the M/M/1 network for the same S and T . In general $Q \leq 1$, and equality holds only for the M/M/1 network itself. The capacity necessary for this

We shall model the delay of our two level system by the following ad-hoc formula

$$T = \frac{1}{C-2S} + \frac{1}{(C-G/2)-eS/2}$$

This equation gives T in terms of C and S , where G is also given in terms of C and S by (28). The first term stands for the repeater-to-station delay, as given by (27) with $U_2=1/2$. The second stands for the terminal-to-repeater delay. It is also based on (27), with the following modifications: Since a repeater cannot listen while talking, the capacity available to each of the terminal groups is $C-G/2$. $S/2$ is the traffic carried by each group, and $1/e$ is the maximum utilization of an infinite population ALOHA.

A three-level organization, as shown in Fig. 5.



WHEN SYSTEM IS DRIVEN TO ITS MAXIMUM UTILIZATION NUMBERS ON LEFT OF LINES SHOW FRACTION OF TIME NODE IS ACTIVE. NUMBERS ON RIGHT SHOW FRACTION OF TIME NODE IS SUCCESSFUL

Figure 5. Structure of the Three-Level ALOHA System.

can improve the system performance even more, for high loads. In the best possible situation, we shall have only two cases of interference: Two messages trying to reach the same repeater will destroy each other; and a message trying to reach a repeater that is itself transmitting will be destroyed without bothering the repeater's transmission. In this case the system can drive the top level to its capacity, and the utilization can be $1/2$. Fig. 6 shows the quality of one-level, two-level and three-level ALOHA systems. For comparison the quality of FDMA with 1024 terminals is also shown.

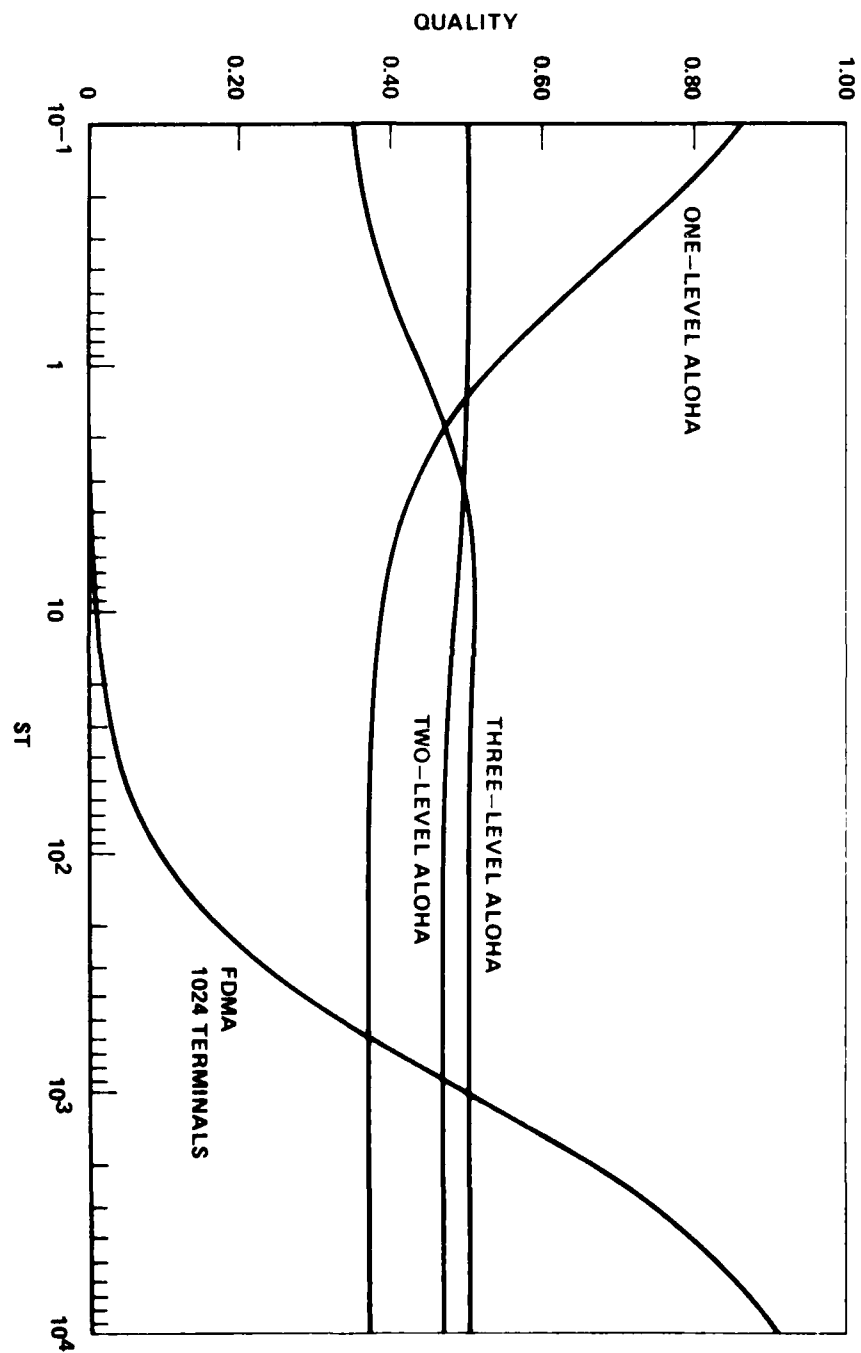


Figure 6. Quality of the One, Two and Three-Level ALOHA Systems.

Four or more levels will never improve the performance of an ALOHA system, as given by our model. To see this consider Fig. 5 again: The numbers on the left of the lines in the two top levels give the traffic per slot that must be offered by the repeaters when the system is driven to its maximum utilization. The numbers on the right give the rate of successful traffic per slot in each hop. In order to get a utilization of $1/2$, each of the top-level repeaters must be active $1/2$ of the time, and will be successful on the average $1/4$ of the time. Each one of the second-level repeaters must be successful $1/8$ of the time, and must therefore be active $1/2$ of the time. The capacity available to each one of the bottom-level infinite population ALOHA systems is $C/2$, and the delay in each will be $T = \frac{1}{C/2 - eS/4}$. When the system is driven to its maximum utilization we have $S = C/2$, and the burstiness of the bottom-level ALOHA system is $\frac{ST}{4} = \frac{1}{4 - e} = .78$. From Fig. 6 we see that at this burstiness a one-level ALOHA system is still better than multi-level systems, and the three-level ALOHA system cannot, therefore, be improved by splitting its bottom level into more levels, even when it is driven to its maximum utilization.

We have just seen that multi-level ALOHA centralized systems can be better than one-level ALOHA when the traffic is heavy, because in the top level we can have a contention system with a small population, which can better utilize its communication resources. Will such a multi-level organization improve networks?

Let us start with one-dimensional networks, and introduce equally spaced repeaters as the top level. We shall have the smallest population of contending repeaters when transmissions go only from one repeater to its two nearest neighbors. Assuming omnidirectional antennas we find that three repeaters, i.e., the source, the destination and its other neighbor, contend at every point. The maximum utilization can therefore go up, from $1/e$ to $4/9$, but the amount of contending traffic has also gone up, from $2NS$ to $3NS$! The reason for the increase in contending traffic is that when we assumed a continuum of terminals and considered a given transmission, the amount of traffic generated exactly at our destination was negligible and our transmission had to contend only with traffic crossing its destination. But when we concentrate the traffic in our repeaters the amount of traffic coming out of a destination is NS , which is not negligible, and must be added to the crossing traffic, equal to $2NS$ as before, in order to get the total contending traffic.

In general, assume each repeater has a range to reach k other repeaters, and, for simplicity, that the distance each message must travel on the repeater-repeater network is a multiple of k . The traffic coming out of each repeater is then NS/k . Each contention system will consist then of $m = 2k + 1$ repeaters and the total traffic in it is $(2k + 1)NS/k = 2mNS/(m - 1)$. Let H be the number of hops necessary, on the average, in the repeater-repeater level. The capacity necessary for this level is therefore

$$C = \frac{1}{U_m} \frac{2mNS}{m-1} + \frac{H}{T} \quad (29)$$

Where U_m is the maximum utilization of an m -repeater ALOHA system. U_m is written explicitly in (26), and substituting we get

$$\frac{1}{U_m} \frac{2mNS}{m-1} = \left(1 - \frac{1}{m}\right)^{-m} 2NS > e 2NS \quad (30)$$

From (30) and (29) we see, even if H is equal to 1, that the repeater-repeater subsystem needs more capacity than the entire one-level network! The detrimental effect of concentrating the traffic and increasing the contention is more important than the gain in the possible utilization of a finite population repeater system. Our conclusion here is, therefore, that if range is no problem, concentrating network traffic into repeaters wastes communication resources. Introducing repeaters can, of course, be an improvement if their range is much larger than the terminals' range, and if this significantly reduces the

number of hops a message must take.

In a heavily loaded two-dimensional ALOHA network we saw that the optimal transmission radius is small. That is, even without repeaters, whenever the traffic is steady we should make our contending terminal system as small and as finite as we dare! Repeaters are not necessary for improving the utilization of a heavily loaded two-dimensional network, and the extra level they introduce is wasteful. Repeaters can be very useful, for networks of intermediate burstiness, if ALOHA is used for terminal-repeater communication and dedicated channels are used for repeater-repeater communication. For a treatment of such mixed-mode networks see [18].

6. Conclusions

Using ALOHA as an access mode for a communication system consisting of a large number of distributed terminals is extremely simple and therefore appealing. But a heavily loaded centralized ALOHA system, in which all messages must reach one common destination, will need e times more bandwidth than the theoretical best (and impossible!) $M/M/1$.

ALOHA networks are in a better position. Since messages have various distributed destinations the channel can be spatially reused: i.e., various transmissions can successfully use the channel at the same time if they are separated spatially and do not interfere at their destinations. The contention between messages is not directly determined by the given traffic, and it can be adjusted by choosing the transmission range.

By modelling a homogeneous and isotropic network by a continuum of terminals we calculated the optimal transmission range. A two-dimensional ALOHA network need be only $\sqrt{e}$ times worse than the corresponding $M/M/1$ network, even when very heavily loaded, as long as the calculated optimal range is not too small to be practical. The calculated range becomes too small when only a few terminals are within range of each other. But the problem of organizing and coordinating a system with a large number of terminals, which was the original motivation for using ALOHA, has disappeared, and other access modes can then be used to advantage, though we have not considered any in this paper.

Since ALOHA networks pay a smaller price for contention than do the centralized ALOHA systems it is harder to improve them by reducing contention. Splitting terminals into power groups can improve any ALOHA system, especially when the traffic is split between groups in a good way, but the resulting improvement in centralized systems is much more significant than the resulting improvement in networks.

In a centralized system all messages must reach the station, and must therefore contend for its ear. A multi-level organization using ALOHA at all levels can improve heavily loaded single-destination systems by having only a small number of intermediate nodes communicate directly with the station. Multi-level ALOHA organizations do not help networks, because choosing the transmission range is a much more effective means for controlling the amount of contention.

References

- [1] Kleinrock, L., "Resource Allocation in Computer Systems and Computer-Communication Networks," *Information Processing 1974, Proceedings of IFIP Congress 74*, Stockholm, August 1974, North Holland, Amsterdam 1974, pp. 11-18.
- [2] Abramson, N., "Packet Switching with Satellites," *AFIPS Conference Proceedings*, 1973 National Computer Conference, Vol. 42, pp. 695-702.
- [3] Lam, S., "Packet Switching in a Multi-Access Broadcast Channel with Application to Satellite Communication in a Computer Network," Computer Systems Modeling and Analysis Group, School of Engineering and Applied Science, University of California, Los Angeles, UCLA-ENG-7429, April 1974. (Also published as a Ph.D. Dissertation, Computer Science Department)
- [4] Kleinrock, L., *Queueing Systems, Vol. II: Computer Applications*, Wiley Interscience, New York, 1976.
- [5] Ferguson, M.J., "A Study of Unslotted ALOHA with Arbitrary Message Lengths," University of Hawaii, Honolulu, Hawaii, Technical Report B75-13, February 1975.
- [6] Roberts, L.G., "ALOHA Packets System with and without Slots and Capture," *Computer Communication Review*, A Quarterly Publication of the ACM Special Interest Group on Data Communication, Vol. 5, No. 2, pp. 28-42, April 1975.
- [7] Akavia G.Y., "Hierarchical Organization of Distributed Packet-Switching Communication Systems," Ph.D. Dissertation, Computer Science Department, University of California, Los Angeles, March 1978.
- [8] Carleial, A. B. and Hellman M. E., "Bistable Behavior of ALOHA-type Systems," *IEEE Transactions on Communications*, Vol. COM-23, pp. 401-410, April 1975.
- [9] Fayolle, G. et al, "Stability and Optimal Control of the Packet Switching Broadcast Channel" *Journal of the ACM*, Vol. 24, pp. 375-386, July 1977.
- [10] Yemini Y. and L. Kleinrock, "On a General Rule for Access Control or, Silence is Golden," in J. L. Grange and M. Gien (Eds.), *Flow Control in Computer Networks*, Proceedings of the International Symposium on Flow Control in Computer Networks, Versailles, 1979, North-Holland, Amsterdam, 1979, pp. 335-347.
- [11] Kleinrock, L., "On Giant Stepping in Packet Radio Networks," Internal Note, UCLA, March 1975.
- [12] Kleinrock L. and J. Silvester, "Optimum Transmission Radii for Packet Radio Networks or Why Six is a Magic Number," *NTC 78, Conference Record of the IEEE National Telecommunication Conference*, Birmingham, Alabama, December 3-6 1978, pp. 4.3.1-4.3.5
- [13] Metzner, J.J., "On Improving Utilization in ALOHA Networks," *IEEE Transactions on Communications*, Vol. COM-24, pp. 447-448, April 1976.

- [14] Gitman, I., "On the Capacity of Slotted ALOHA Networks and Some Design Problems," *IEEE Transactions on Communications*, Vol. COM-23, pp. 305-317, March 1975.
- [15] Kleinrock, L., "Performance of Distributed Multi-Access Computer Communication Systems," *Information Processing 1977, Proceedings of IFIP Congress 77*, Toronto, August 1977, North Holland, Amsterdam, 1977, pp. 547-552.
- [16] Ferguson, M.J., "A Bound and Approximation of Delay Distribution for Fixed-Length Packets in an Unslotted ALOHA Channel and a Comparison with Time Division Multiplexing (TDM)," *IEEE Transactions on Communications*, Vol. COM-25, pp. 136-139, January 1977.
- [17] Lam, S., "A New Measure for Characterizing Data Traffic," *IEEE Transactions on Communications*, Vol. COM-26, pp. 137-140, January 1978.
- [18] Akavia, G.Y., and L. Kleinrock, "On the Advantage of Mixing ALOHA and Dedicated Channels," submitted for publication.

Hierarchical Use of Dedicated Channels

Abstract

We consider efficient organizations for communication resources which are accessed by a large number of geographically distributed terminals. Developing a model for systems built with dedicated channels, we answer the following questions: What is the role of hierarchies in organizing large communication nets? How should a large network be decomposed into smaller parts? What cost versus performance gains can be achieved by such a decomposition?

Assuming that performance is specified and that the goal is to minimize the necessary cost, we define *quality* and *burstiness* and find the following: Dedicating channels is reasonable when the traffic is steady (i.e., not bursty), but when the traffic is bursty the cost of simple dedicated-channel systems grows too fast with the number of terminals. By introducing *regular* hierarchical structures we show that the cost of bursty systems can be significantly reduced. The *optimal* structure must be *balanced*, and the ratio of the contribution of the different levels to both cost and delay is simply determined by a few key system parameters.

We consider two technologies: line and broadcast. The cost of the best bursty *line* system grows with the dimensionality of the space in which terminals are distributed. The cost of the best bursty *broadcast* system is similar to the cost of *one* dimensional line systems and is independent of dimensionality. It follows that bursty broadcast systems have an advantage over line systems in two or more dimensions.

The above apply to both *centralized* systems, in which messages originate in the distributed terminals but are directed to one common destination, and to *networks*, in which both sources and destinations of messages are distributed.

Organizing a two-dimensional network imposes a tessellation on the plane. We compare the three regular tessellations and analyze the relevant tradeoffs. When using the best number of levels, as a function of burstiness, tessellating the plane with hexagonal tiles (and forming a triangular network of communication lines) is usually optimal.

1. Introduction

Designing a communication network for a given traffic requirement consists of balancing cost and performance. Faced with the task of analyzing networks, we must abstract the relevant features of traffic, performance and cost in order to arrive at a manageable model. In this paper we develop such a model and use it to answer the following questions: What is the role of hierarchies in organizing large communication nets? How should a large network be decomposed into smaller parts? What cost versus performance gains can be achieved by such a decomposition? To motivate the abstractions necessary to arrive at our model consider the following simple example:

Assume messages originate at m different sources (buffered terminals). Assume that the appearance of messages at each source is a Poisson process with rate S/m messages per second, and that the length of messages has an exponential distribution. Let us choose the information unit so that the average length of a message is equal to 1; this is simply a convenient normalization, which is equivalent to measuring communication capacity in messages (of an average length) per second, instead of measuring in bits per second. Assume all messages are directed to one destination (computer), which we shall sometimes call the station.

Consider the two cases shown in Fig. 1.

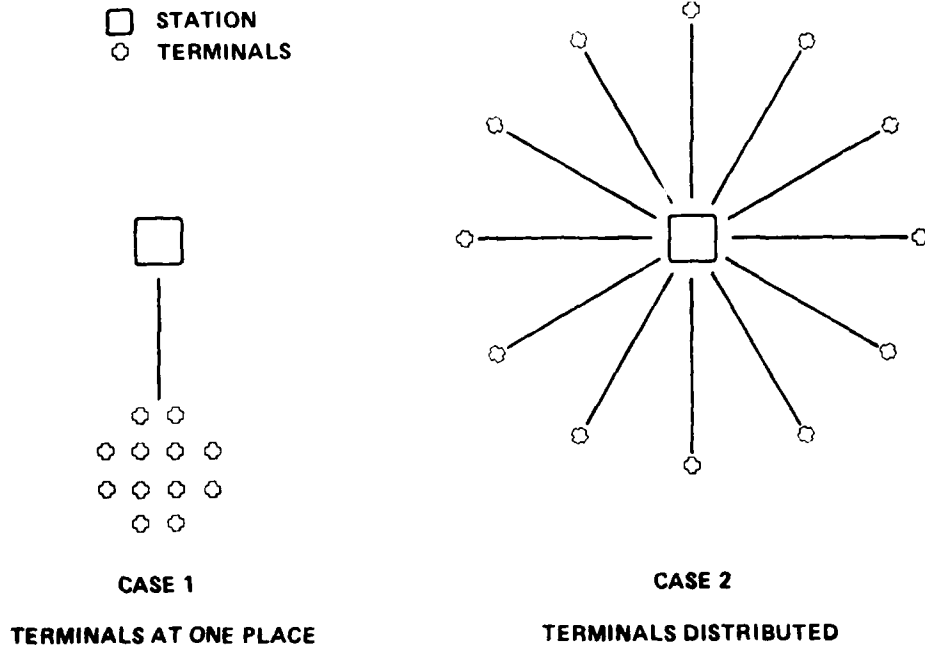


Figure 1. Centralized versus Distributed Terminals.

In both cases all terminals are at the same fixed distance from the station. In case 1 all terminals are at one and the same location. They can, therefore, share a single communication channel. In case 2 the terminals are spread out around the station, and we shall connect each one to the station by a separate, individual channel.

How should we compare these communication systems? Having fixed the structure of both systems, and since the distances from all terminals to the station is the same in both cases, we shall ignore for the moment the question of distances and cost, and shall characterize both systems by the relation between the following three parameters:

- S Total rate of messages transmitted (messages per second)
- T Average total time spent by a message in the system (seconds)
- C Sum of the capacities of all communication resources used (messages per second)

In order to compare the two systems of Fig. 1 let us first find the relation between S , T and C that characterizes each of them.

In case 1 all sources are in one place and are connected to the destination by a single communication channel. Each message will join a queue at the terminal end of the channel, and when its turn comes, will be transmitted to the destination. We thus have a classical $M/M/1$ queueing system [1] with arrival rate S and service rate C (messages per second). The average total time T a message

spends in the system (in queue and in service) is given by

$$T = \frac{1}{C - S} \quad (1)$$

In case 2 each terminal is connected to the station by an individual channel. If C is the total capacity available, let us connect each source to the destination by a channel whose capacity is C/m . Each message will therefore have to pass through one of m identical queueing systems (with arrival rate S/m and service rate C/m each.) The relation between capacity and average time in this system is simply

$$T = \frac{1}{C/m - S/m} = \frac{m}{C - S} \quad (2)$$

If the communication capacity we use is predetermined, it is natural to compare the delay in the alternative organizations. In our case, for a given C and S , let T_1 and T_2 be the time spent in case 1 and case 2 of Fig. 1 respectively. Forming the ratio of (1) and (2) we get

$$\frac{T_1}{T_2} = \frac{1}{m} \quad (3)$$

The M/M/1 system of case 1 is, with the given assumptions on the statistical nature of message arrival and length, the best we can achieve, i.e., we pay the only the unavoidable price for *queueing*, and nothing more. In case 2 we have the same queueing effect, but in addition pay a significant amount for the decision to dedicate a part of the channel to each of the terminals. Equation (3) says that a system with m dedicated channels is m times worse than sharing one M/M/1 channel! For this and other scaling results see [2].

In our simple example, T_1/T_2 does not depend on either S or C . But even in the general case, the ratio of times used to compare two systems is a dimensionless number. It can, therefore, depend on S and C only via their dimensionless ratio S/C , which is the utilization of the communication channel, usually denoted in queueing literature by ρ . When $S \ll C$, we say that the system is *lightly loaded*. When S is very near C , we say that the system is *heavily loaded*. When $S \geq C$ the system is *overloaded and unstable*; we shall not treat this case. Both (1) and (2) give the average delay in the steady state of a stable system.

Equation (3) compares M/M/1 and the dedicated channels scheme when C and S are given. How do they compare if T and S are given and we want to minimize the necessary capacity? Let C_1 and C_2 be the capacities necessary in cases 1 and 2. Inverting (1) and (2) and forming the ratio we get

$$\frac{C_1}{C_2} = \frac{ST + 1}{ST + m} \quad (4)$$

It is not surprising that the dimensionless ratio given in (4) depends on S and T only via their dimensionless product ST . We shall call the inverse of ST the *burstiness* [3] of the system. When ST is small ($ST \ll 1$) the system is *bursty*. When ST is large ($ST \gg 1$), the system is *steady*. When the traffic is bursty there are only a few messages in the system. There is little congestion, and the delay suffered by messages is mainly determined by the time necessary to transmit them. The communication resource is only lightly utilized in a bursty system. When the traffic is steady the communication resource is heavily utilized and the delay is mainly determined by the congestion.

Definitions equivalent to our burstiness were introduced independently by others [4,5]. This is not surprising, since ST is the only dimensionless number one can form with S and T . Lightly loaded systems are bursty, and heavily loaded systems are steady, so we shall sometimes use these terms interchangeably. But we shall use the terms bursty and steady when we wish to stress the fact that S and T are given, and that C is to be determined in the design process. We shall also use the terms bursty and steady to describe the traffic a given system has to carry.

Equations (3) and (4) may look very different intuitively, even though they compare the same pair of systems. If we assume S and T are given and compare the needed capacity we see that C_1/C_2 depends on burstiness: when the system is very bursty ($ST \rightarrow 0$) dedicated channels are m times as bad as $M/M/1$, when the system is very steady ($ST \rightarrow \infty$) dedicated channels are almost as good as $M/M/1$. But if we assume that C is given and compare the delay as a function of load then (3) tells us that dedicated channels are *always* m times as bad as $M/M/1$. Which comparison is more meaningful?

In a real commercial environment we may be constrained to use a communication channel with one of several predetermined capacities. Comparing delay will then be the right tool for evaluating alternative system organizations, and (3) will be more meaningful.

However, for the purpose of this paper, we shall assume that capacity can be freely chosen in the course of a system design. The client of the design will specify traffic and performance, and we shall evaluate different designs by the resources necessary in each of them. While this attitude ignores some of the real-life constraints, we feel it gives a much better understanding of many important technical issues.

2. Designing Distributed Communication Systems

Why is it that the terminals in case 2 of Fig. 1 cannot form one queue and use one common channel? One may say that the terminals are *distributed in space*, and therefore cannot share a channel. This statement is reasonable if we are committed to using lines for communication, but in general it should and it can be made more precise. While lines connect pairs of points, other communication technologies have the *broadcast* property: a transmission made by one terminal will be heard by all others. Consider the following gedanken experiment: Assume our terminals have a strong empathy and that, as a result, each one of them senses, immediately and with no error, the fact that another becomes ready to transmit. Despite being distributed in space such a set of terminals can easily form one queue and share one broadcast channel. We may say that if perfect knowledge of who is ready to transmit was available, then being distributed in space would have been of no consequence.

Consider now another gedanken experiment: There is no empathy between terminals, but there is a demon who has perfect knowledge of who is ready to transmit. Assume also that terminals will transmit only when instructed to do so by the demon, and that these instructions arrive free and without delay. Then, once again, the terminals can easily share a broadcast channel: a queue will form in the demon's head, and the demon will instruct the terminal at the head of the queue to transmit. We see from this hypothetical example that it is enough to have perfect information in one place, if that one place could perfectly control all transmissions.

The problem of real distributed communication systems is that the control of transmissions is distributed, and must be based on distributed information. The information that is available at each place is therefore partial and old. We have no perfect empathy and no cooperative demon. Faced with this reality people have developed many schemes for deciding which terminal will use which part of the communication resources at a given time. These schemes, often called *access modes*, usually utilize some of the following ideas: central control using preallocation (TDMA, FDMA) or polling [6], reservations [7,8,9], ALOHA [10], and carrier sense [11].

It would have been nice to be able to completely characterize all possible access modes, and say which one is best for which range of system parameters. But we are far from achieving such a goal. We know no complete characterization of access modes. The performance of many of the known access modes is extremely hard to obtain in an analytic way because they involve complex systems of interacting queues. While it is often easy to evaluate an access mode for a small range of parameters by simulation, it is hard to use simulation to get insight as to which access mode is best for which range of

parameters.

Rather than trying to treat the ensemble of all possible access modes we shall concentrate on one of the simplest - using *dedicated* channels. This is reasonably good when the traffic is steady, but bad when the traffic is bursty. We shall assume the communication system has very many terminals, distributed over very large distances, and ask: For a given traffic and required performance, can the cost of a very bursty system be reduced by a *hierarchical organization*? Before trying to answer this question, let us say how we shall describe traffic, specify performance, and calculate cost.

To specify traffic we shall assume m , the number of terminals, is very large, that terminals are uniformly distributed in their geographic region, and that all terminals contribute equally to the traffic. The reason is that we are interested in hierarchies that arise in the design process, and not in hierarchies that are imposed by the topology and traffic requirements. It is also often true that the uniform case is the worst case for a distributed system: if traffic was especially concentrated in some terminals or regions then the system would be less distributed. In addition we shall assume different messages appear independently. When we treat very bursty traffic the exact distribution of message interarrivals is irrelevant, and only S , that total rate of messages, will appear in our formulas.

Delay will be our only performance measure, and we shall ignore the very important issue of reliability. Indeed, only the *average* delay T will appear in our formulas, but essentially all results will remain valid when the variance, range or distribution of acceptable delay values is specified in addition to the average delay. Meister et al [12] propose and analyze a performance measure that can influence the variance of delay. We shall show later that we can achieve equivalent results by adjusting our cost measure.

The cost of communication depends on technology. We shall classify the very many technologies possible into two groups: *line* systems and *broadcast* systems; and shall assume a cost measure for each group.

A line enables the two points at its ends to communicate. The line can be a tight string, a pair of wires, a coaxial cable, or a light guiding optical fiber. Line-based systems have many advantages, but depend, of course, on a line arriving at every point that needs to communicate. We shall assume that the cost of a line system consists only of the cost of lines, and that the cost of a line channel is directly proportional to the a -th power of its length, and to the b -th power of its capacity. By choosing $b < 1$ we model the economy of scale usually present when building or buying a large capacity channel. When $a < 1$ we actually can take into account the cost of equipment at the ends of the line, which we do not consider explicitly.

The second type of communication technology we shall deal with is that of broadcast systems. The main property of broadcast channels is, that for better or worse, everybody within range can talk, listen and interfere with everybody else; that is, they all hear every transmission. When everybody is within range of everybody else we have a one hop system - every message can arrive from source to destination in one hop. If the transmission range is less than the distance spanned by the terminals we have a multi-hop system. A message may have to be transmitted more than once, at first from its source and then from intermediate 'relays', in order to arrive at its destination. In a multi-hop system it is possible for two different transmissions to successfully use the same broadcast channel at the same time, if they are not within range of each other, i.e., a *broadcast channel can be spatially reused*. When choosing a transmission range we must, therefore, face the following tradeoff: If we choose a large range we shall need few hops, but will cause a lot of interference and monopolize the channel in a large region. We analyze this tradeoff, but ignore the following fact: Range is determined by transmission power, among other factors, and power is seriously limited when terminals are mobile.

When dealing with broadcast systems we shall entirely ignore the cost of equipment (transmitter, receiver, antenna, power source) and consider only the amount of broadcast *bandwidth* used as the cost of the system. The motivation is that technology will make the equipment cheaper and cheaper, but that the bandwidth is now and is likely to remain a truly scarce resource, especially as the overall communication traffic grows. We shall assume the cost of a dedicated broadcast channel with capacity C is given by C^b and ignore a technology-dependent multiplicative constant. Usually b will be smaller than one: there is some cost in bandwidth when a separate channel is created, and wide band channels are therefore relatively cheaper.

The division of all possible communication systems into either line or broadcast systems is, of course, somewhat arbitrary. On the one hand, a broadcast transmitter with a directional antenna and beam can become part of a line system, as the microwave links of the telephone system show. On the other hand, a broadcast system like ALOHA can be implemented on a set of lines [13]. Communication satellites, a prime example of broadcast technology, are actually used by the international telephone community as 'lines', i.e., for point-to-point communication connecting a single source with a single destination. We consider both this division into lines and broadcast systems, and the cost assignments we made, to be useful abstractions, that help isolate the issue of being distributed, which is our main interest here.

Real systems are built slowly. Investments have to be based on estimates of future demand, and the demand in the future is influenced by the existence of the system and the quality of service. We shall ignore this interaction over time, and assume our systems are built in order to satisfy the known demand and service requirements at a given time.

3. Decomposition and Resource Allocation

Having specified our performance and cost measures, let us return to our m equally talkative terminals, all of whom wish to communicate with the single station. Denote by L the 'typical' linear dimension of the region over which terminals are distributed, and assume a line-based communication system is built to connect all terminals to the one station. Since we assume that the cost of every line is proportional to the a -th power of its length the total cost of our centralized system must be proportional to L^a . The total cost must also be proportional to the b -th power of the typical line capacity. When the traffic is very bursty the typical capacity must be $1/T$ (see equation (1)), and it follows that the total cost is proportional to $1/T^b$. The total cost D can therefore be written, without loss of generality, as

$$D = \frac{L^a}{T^b} f \quad (5)$$

Given our assumption on the cost of individual lines, the dependence of D on L and on T is an inevitable result of the traffic requirements, i.e., of wanting to communicate (across distances that are typically L) over lines (whose capacity must typically be $1/T$.) The f appearing in (5) shows how the system cost depends on its being distributed. f contains some geometric constants, and a dependence on m , the number of terminals. We shall usually ignore the constants, and address the dependence on m : How fast does f grow with m ? Must it grow that fast?

Assume we have a procedure for designing a very bursty centralized communication system, given m , the number of equally talkative and uniformly distributed terminals. Such a design procedure can be completely characterized by its f -function, defined by (5).

Applying a given design procedure to a communication system with very many terminals may be too expensive. Can we reduce cost by decomposing the system into small parts, and by applying the given design procedure to each part separately? How should we decompose a large system and how should we allocate resources to the different subsystems? We shall start with the latter question. Assume the cost of the j -th subsystem is given by (5), i.e.,

$$D_j = \frac{L_j^a}{T_j^b} f_j$$

and that the total system cost is $D = \sum_j D_j$. Assume that the delay measure T is given by the following weighted average

$$T = \sum_j S_j T_j / S \quad (6)$$

where S_j is the traffic carried by the j -th subsystem and S is the total traffic. If we now choose the T_j in order to minimize D given T (or in order to minimize T given D) we get the following cost:

$$D = \frac{1}{(ST)^b} B^{b+1} \quad (7)$$

$$\text{where } B = \sum_j \left(L_j^a S_j^b f_j \right)^{1/(b+1)}$$

Minimizing the cost of a hierarchical structure often involves minimizing B given in (7), which we shall call the B -term

When resources are allocated to subsystems in the optimal way, which leads to (7), we also get

$$\frac{D_j}{D_k} = \frac{S_j T_j}{S_k T_k} = \left(\frac{L_j^a S_j^b f_j}{L_k^a S_k^b f_k} \right)^{1/(b+1)} \quad (8)$$

That is, the contributions of subsystems to the delay measure and to the cost are directly proportional to their contribution to the B -term.

When our subsystems consist of a single line each Equation (7) is very similar to Kleinrock's optimal capacity assignment [16], with the following difference: by restricting ourselves to very bursty traffic we can handle cost functions with any b , not just the $b=1$ case. When the traffic is very bursty there is also a simple equivalence between modifying the delay measure to $T^{(k)}$ of Meister et al [12] and modifying the cost measure by substituting b/k for b .

When writing (6) we have assumed that the routing of individual messages does not depend on the state of the network, i.e., routing is not adaptive. We see that no matter what b is, the B -term is a concave function of S_j and the best routing must therefore result in a tree-like network - it does not pay to split the traffic from a given source to a given destination and route each portion differently.

When the performance measure specified includes the distribution of delay values, equation (6) may be too strict, since it imposes a similar distribution on every one of the subsystems. Equation (6) can then be considered a heuristic, and the resulting allocation may be suboptimal.

4. Regular Hierarchical Structures

Having decomposed a communication system, equation (7) gives a way to allocate resources to its various parts. We do not know which is the optimal way to decompose a large system for our goal of minimizing cost, so we shall use another heuristic. To introduce it, consider the following two-level structure. Assume the m terminals are uniformly distributed in a region of n -dimensional space, and divide this region into P congruent regions. Place a concentrator in the middle of each region, connect all P concentrators to the station according to a given design procedure, and connect all terminals in a given subregion to 'their' concentrator according to the same design procedure. For simplicity of our formulas we shall assume that all subregions have the same shape as the original region, and will ignore the constant coefficients that depend on this common shape and on the dimensionality.

We shall call this hierarchical system a two-level *regular* hierarchical system where the word *regular* refers to the fact that all regions are of the same size and shape, and that all concentrators are placed in the middle of their regions. We shall call the communication subsystem connecting concentrators to the station the *top* level, and the subsystem connecting terminals to concentrators the *bottom* level. The top level consists of a network with the P concentrators acting as terminals, and the bottom level consists of P networks with m/P terminals each.

Let L be the typical linear size of the original n -dimensional region. The typical linear size of each one of the P subregions is $L(1/P)^{1/n}$, and the total traffic arriving at each concentrator is S/P . Applying (7) to both levels we find that the contribution of the bottom level to the B -term is

$$P \left[L^{1/n} (1/P)^{1/n} (S/P)^b f(m/P) \right]^{1/(b+1)}$$

Where we have shown explicitly the dependence of f on m/P , the number of terminals in every subregion. The contribution of the top level to the B -term is

$$\left[L^a S^b f(P) \right]^{1/(b+1)}$$

Adding gives the B -term of the two-level regular hierarchical system:

$$B = \left[L^a S^b \right]^{1/(b+1)} \left[f(P)^{1/(b+1)} + P^{(1-a)/(b+1)} f(m/P)^{1/(b+1)} \right] \quad (9)$$

Which P will give the least cost two-level system? Are two levels better than one? The answer to the second question will follow from the answer to the first, since when $P=1$ or $P=m$ the two-level system reduces to a one-level system. This is reflected in (9) since $f(1)=0$: when we have to connect one terminal, which is 'uniformly' distributed over its region, to a station in the middle of the region there is nothing to do, and no cost is incurred.

To find the best P that will minimize B we must say something about the f -function. For simplicity assume that when m is large the following is a good approximation:

$$f(m) = m^g \quad (10)$$

Assuming that P satisfies $m \gg P \gg 1$, so that both P and m/P are large, we can substitute (10) into (9) and get

$$B = \left[L^a S^b \right]^{1/(b+1)} \left[P^{g/(b+1)} + P^{(1-a)/(b+1)} (m/P)^{g/(b+1)} \right] \quad (11)$$

Differentiating B with respect to P we see that $dB/dP=0$ when

$$g^{b+1} P^g = \left[g - 1 + a/n \right]^{b+1} (m/P)^g P^{1-a/n} \quad (12)$$

Substituting the P determined by (12) into (11) we see that the cost of the two-level structure, optimized with respect to P , is proportional to m^h , where

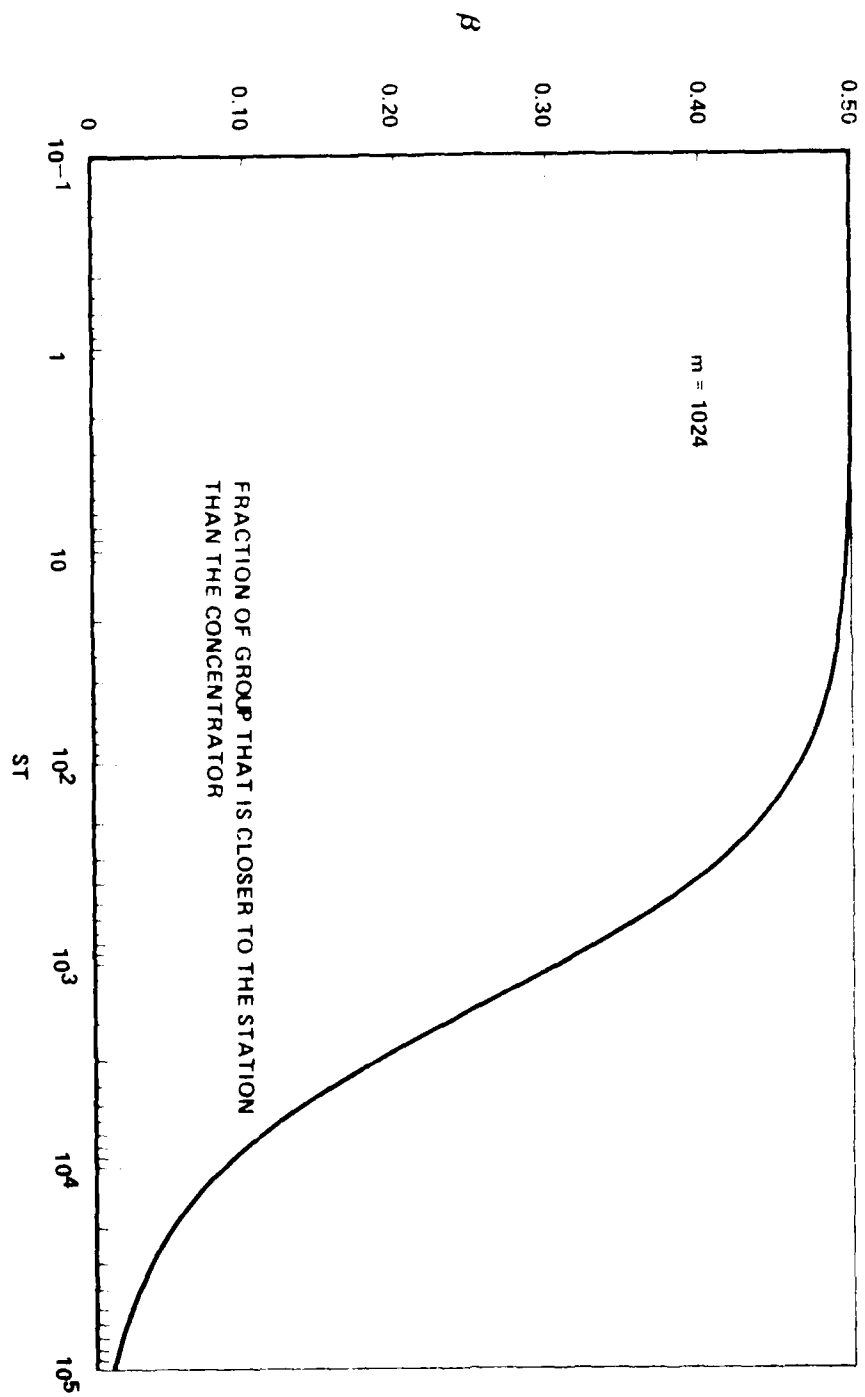


Figure 7. Optimal Placement of Concentrators.

Substituting $\beta=0$ in (29) and searching for the best P_1 and P_2 we see that the cost of the best two-level go-forward system is

$$D = \frac{L}{4} S + \frac{8}{9} \frac{L}{T} \sqrt{2m} \quad (30)$$

Comparing (30) to (25) we see that, for large m , the two-level go-forward system is better than the one-level system for all values of ST . But when the system is very steady, there is very little to gain by introducing a two-level structure.

Fig. 7 shows the optimal β as a function of ST . When the traffic is bursty we should use regular systems ($\beta=1/2$) and as ST grows β becomes smaller, and the best systems with a very steady traffic are go-forward systems ($\beta=0$). But Fig. 8 shows that the idea of choosing the best place for the concentrators as a function of ST is almost irrelevant! Fig. 8 shows the cost of the two-level regular system, the two-level go-forward system and the one-level system as a function of ST . The costs were normalized, for each value of ST , by the cost of the two-level system with the best concentrator placement for that ST , as given by (29) when β is chosen to minimize D . Assume we have to design a system with a given ST , and consider the following decision: we shall use the regular two-level system, with the optimal number of groups for the given ST , as long as it is better than the one-level system. Otherwise we shall simply use the one-level system. From Fig. 8 we see that if we follow this procedure, instead of trying to find the two-level system with the optimal routing policy, then our expenses will be larger by at most 1% ! A similar conclusion applies to networks [3]: if the one-level system is not good enough we may consider only *regular* multi-level systems, and lose almost nothing.

8. Distributed Dedicated-Line Networks

Until now we have only dealt with the *centralized* system case. That is, the sources of messages were distributed, but all messages were directed to one destination, i.e., the station. We shall now begin treating the case of communication systems with distributed destinations, which we call *networks*. When analyzing networks we shall be able to use many of the results obtained for centralized systems. To see how, consider first one-dimensional networks built with dedicated line channels.

Assume terminals are located at fixed intervals along our one-dimensional networks, and let l be the distance between any pair of nearest neighbors. Each terminal wishes to communicate with all other terminals. The traffic of messages between any two terminals is a Poisson process, whose rate depends only on the distance between terminals, and not on their identity. That is, all terminals are identical in their statistical properties. We need the distribution of distances travelled in order to completely specify the traffic. However, most of our results will depend only on N , the *average* distance travelled.

Let us assume that our network is 'infinite', i.e., its total size is so much larger than N that an insignificant fraction of terminals are affected by the boundaries of the network. It makes no sense to talk about the total traffic carried, so let S_u denote the traffic coming out of a *unit length* of the network. D_u will similarly denote the budget invested in a *unit length* of the network.

Our motivation for choosing an entirely uniform universe may now be restated: If some terminal had an especially large communication requirement, or if it was especially central in some sense, we would naturally treat it in a special way when designing a good system. We, however, are interested in the differentiation between terminals that appears when hierarchies are built in an entirely uniform environment, even though no terminal is special to begin with.

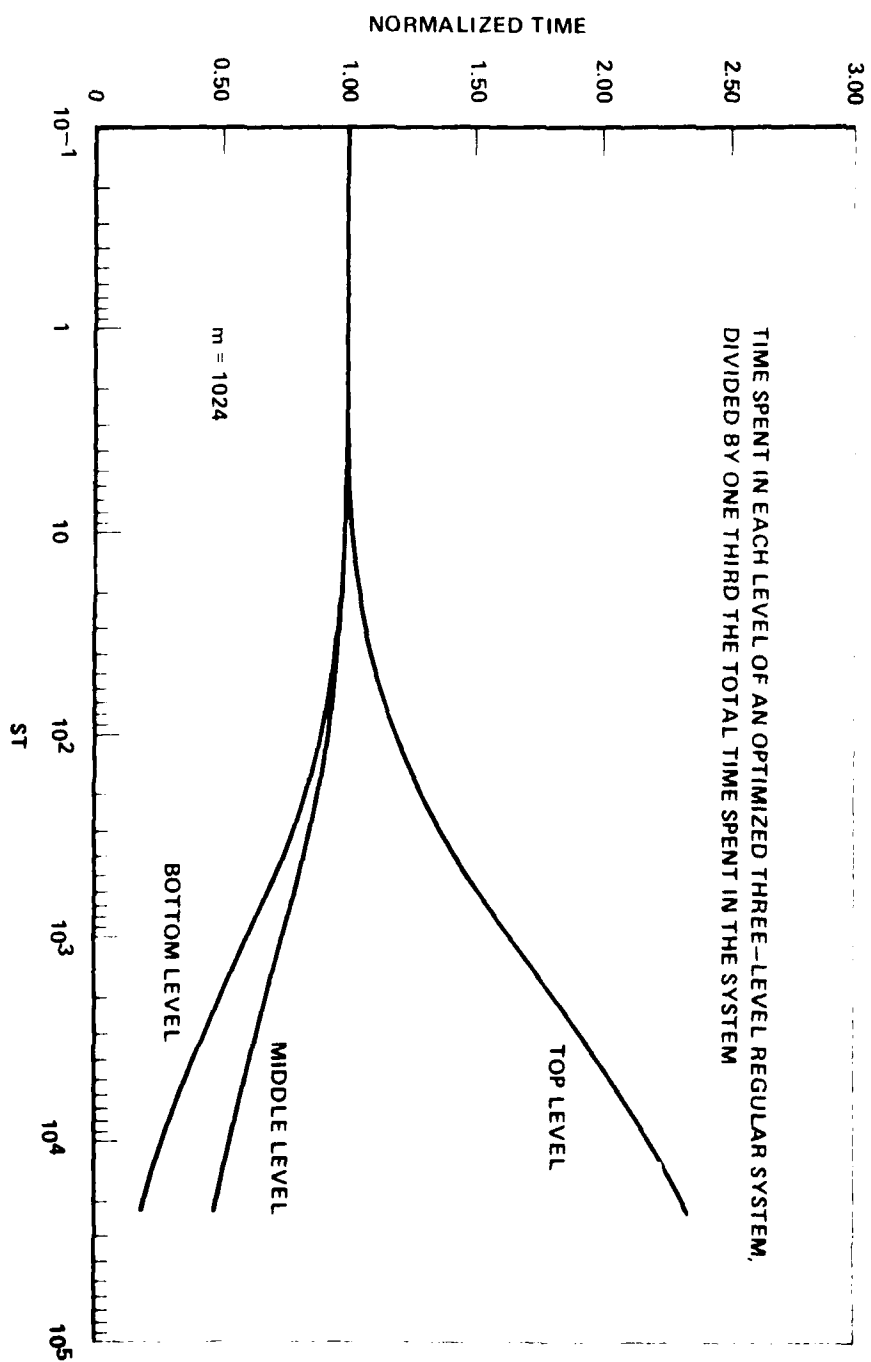


Figure 5. The Time Spent in the Three Levels.

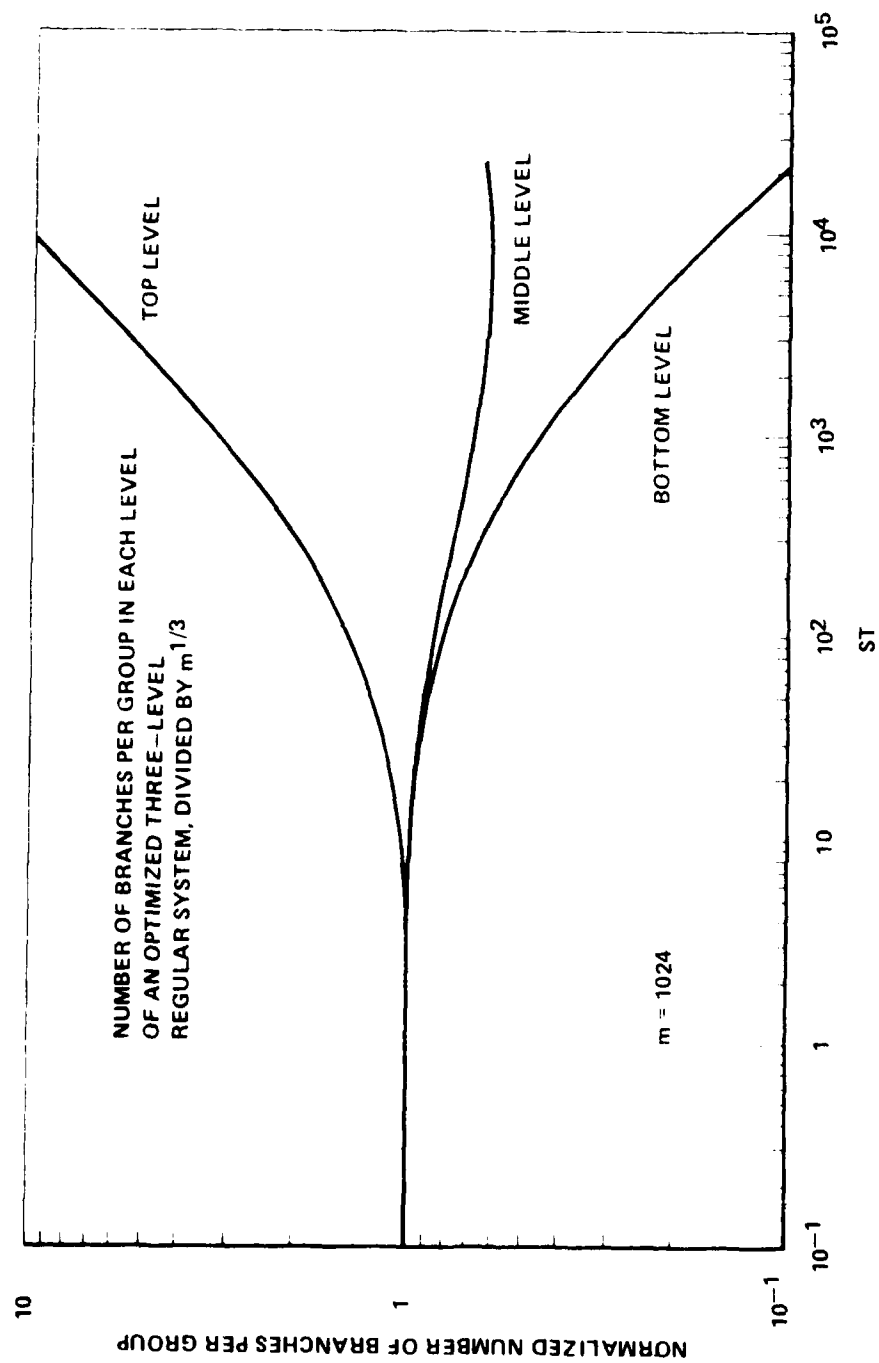


Figure 6. The Number of Branches in the Three Levels.

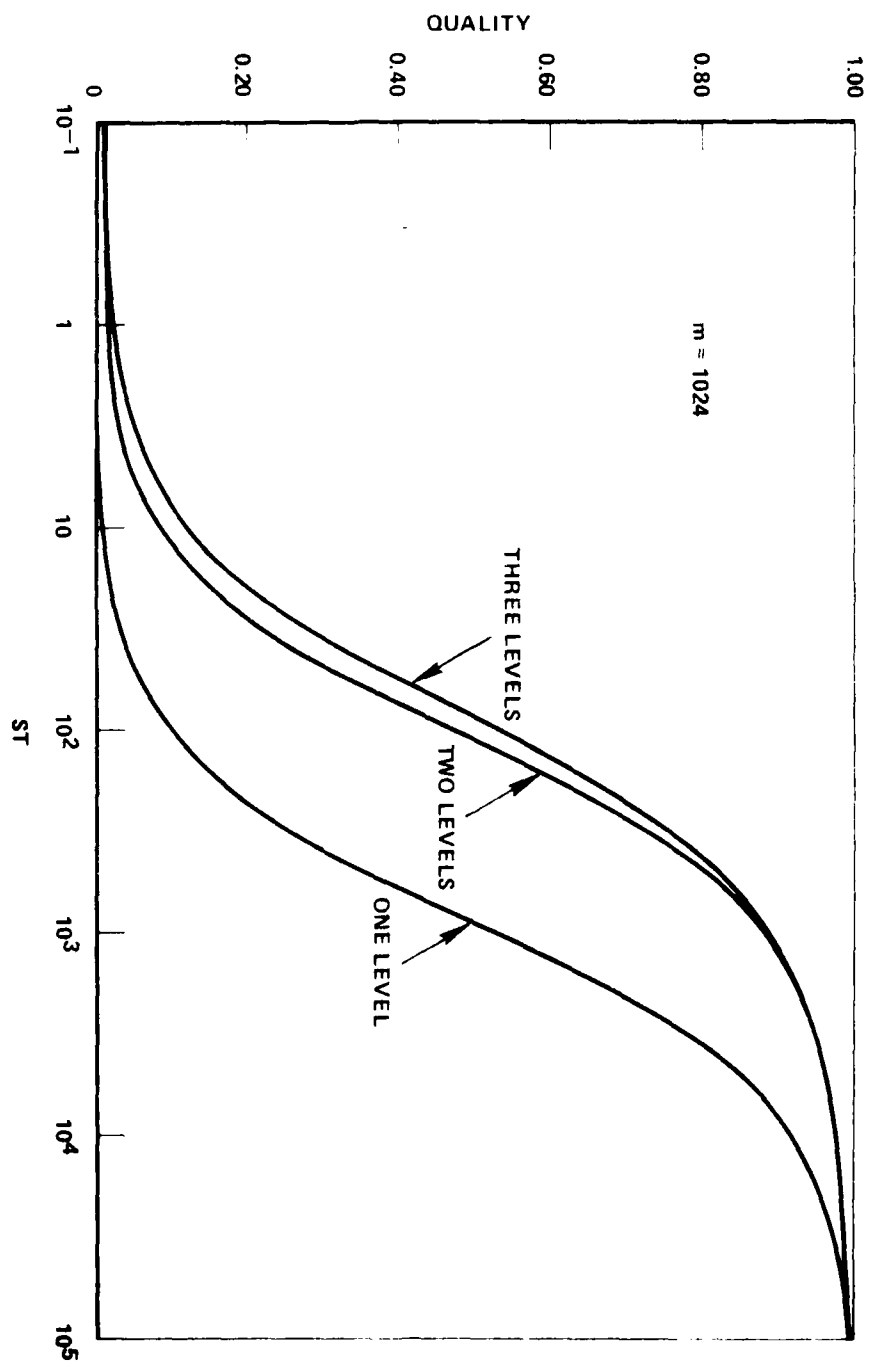


Figure 4. Quality of Non-Bursty Regular Systems.

Fig. 4 shows the quality of the one-level, two-level and three-level regular systems for 1024 terminals. When the traffic is very bursty the three-level organization is better. When ST grows its advantage becomes less pronounced, and if the curves in that figure were drawn fine enough we could have seen that the two-level system and then the one-level system takes over. Fig. 5 shows, for a three level organization, the ratio between the time spent in every level and one third of the total time spent in the system. Fig. 6 shows, for a three level organization, the ratio between the number of branches in every level and $m^{1/3}$. In both of these last two figures, the convergence of all three curves to a common point when $ST \rightarrow 0$ is a manifestation of the balanced nature of bursty systems under optimal capacity assignment.

Multi-level regular systems are much better than the one-level system when the traffic is bursty. Why do they become progressively worse than the one-level system as the traffic becomes steadier?

In the regular systems the concentrators are placed in the middle of their group. This means that some messages will take a route which is longer than the direct distance from their origin to the station. When the traffic is bursty, this effect is negligible compared with the gains resulting from sharing the long lines. But when the traffic is steady, sharing leads only to a small gain, and the extra distance travelled is significant. When ST is very large, we see by comparing (24) and (26) that the two-level regular system costs $LS/(4 P_1)$ more than the one-level system. This extra cost is a direct expression of the extra distance travelled. Half the terminals, i.e., those terminals whose concentrator is further away from the station than they are, will have to travel an extra distance equal to twice the distance to their concentrator. The average extra distance travelled is therefore simply the average terminal-concentrator distance, which is equal to $L/(4 P_1)$.

We can decrease the extra distance travelled by placing the concentrators nearer to the station. Let us, for simplicity, adopt the policy that all concentrators will be placed so that a fraction β of their group will be on the side near the station. In analogy with (26) we get that the cost of the two-level system built with this policy is

$$D = \frac{L}{4} S (1 + 4\beta^2 P_1) + \frac{L}{7} B^2 \quad (29)$$

where

$$B = \frac{1}{3} \sqrt{2P_1} + \frac{2}{3} \left[\beta^{3/2} + (1 - \beta)^{3/2} \right] \sqrt{P_2}$$

When $\beta = 1/2$ this equation reduces, of course, to (26).

For a given value of ST , which P_1 , P_2 and β will give the least cost system? For a given β , finding the best P_1 and P_2 is easy, and the best β can then be found numerically. As is intuitively clear, for bursty traffic the best β is equal to $1/2$. When the traffic becomes steadier the best β becomes smaller, and when the traffic is extremely steady the best β is equal to zero.

It is interesting to note that, for any given β , the system with the optimal group size obeys a balance principle: The excess budget is invested equally in the two levels and the average delay in the two levels is the same.

When $\beta = 0$ the system has a nice property that we formalize thus: A communication system in which the length of the route taken by any message is equal to the direct distance from its source to destination will be called a *go-forward* system.

Let D_1 be the amount of money invested at the concentrator station level (the top level) and D_2 be the amount invested at the terminal-concentrator level (the bottom level). Let T_1 and T_2 be the average time a message spends in the top level and the bottom level respectively. From (23) and (27) we get the following

$$\frac{D_1 - LS/4}{D_2 - LS/(4P_1)} = \frac{T_1}{T_2} = \sqrt{1 + \frac{9}{8} \frac{ST}{m}} \quad (28)$$

The first equality is not specific to regular systems. It follows directly from (22) that whenever we consider two sets of lines in a communication system with an optimal capacity assignment, the ratio of their contribution to the average delay is equal to the ratio of excess the budget invested in them. The second equality sign shows how both of these ratios depend, in a two-level regular system, on ST .

When $ST \rightarrow 0$, (27) shows that $P_1 = P_2$ and (28) is then just a specific case of (8); every regular two-level system must be balanced when bursty. When the system becomes less bursty P_1/P_2 , T_1/T_2 and D_1/D_2 grow. There are more branches than leaves per branch, more of the budget is invested in the top level, and the message spends more time in the top level. When ST becomes large enough, i.e., the system becomes very steady, we get from (27) that P_2 is less than one! This means that for large enough ST a one-level system will be better than a two-level system. Our optimized two-level system is trying to achieve the one-level performance by 'eliminating' the unnecessary bottom level, or at least by lessening its effect.

r -level regular systems can be optimized by applying (27) and (28) to every two consecutive levels. As an example, let us solve the three-level case.

A three-level regular system will have P_1 branches at the stem, each of which splits into P_2 twigs, each of which carries P_3 leaves. The two top levels can be considered as a two-level regular system with $P_1 P_2$ terminals. The two bottom levels can be considered as a set of P_1 identical two-level regular systems with $P_2 P_3$ terminals, each with a total throughput S/P_1 . P_1, P_2 and P_3 must satisfy

$$P_1 P_2 P_3 = m$$

Applying (27) to the two top levels and to the two bottom levels we get

$$\begin{aligned} \frac{P_1}{P_2} &= 1 + \frac{9}{8} \frac{S(T_1 + T_2)}{P_1 P_2} \\ \frac{P_2}{P_3} &= 1 + \frac{9}{8} \frac{S(T_2 + T_3)}{P_1 P_2 P_3} \end{aligned}$$

where T_1 , T_2 and T_3 are the average times spent in the top, middle and bottom level, correspondingly, and they satisfy

$$T_1 + T_2 + T_3 = T$$

Applying (28) to the two subsystems, we get

$$\begin{aligned} \frac{T_1}{T_2} &= \left(\frac{P_1}{P_2} \right)^{1/2} \\ \frac{T_2}{T_3} &= \left(\frac{P_2}{P_3} \right)^{1/2} \end{aligned}$$

We therefore have six equations for six unknowns. While we do not have an analytic solution for them, a numerical one is easy to obtain.

$$C_i = S_i + \frac{D_e}{L_i^a} \frac{\sqrt{S_i L_i^a}}{\sum_i \sqrt{S_i L_i^a}} \quad (22)$$

where

$$D = \sum_i S_i L_i^a + \frac{1}{ST} \left(\sum_i \sqrt{S_i L_i^a} \right)^2 \quad (23)$$

and

$$D_e = \frac{1}{ST} \left(\sum_i \sqrt{S_i L_i^a} \right)^2$$

A certain minimum budget is necessary for carrying the given traffic, even if we are willing to tolerate a very large T . D_e is the *excess budget*, invested in order to make the delay finite.

We shall now consider in detail the case of one-dimensional centralized systems, in which the cost of a line is directly proportional to its length, (i.e., $a=1$). Let our m terminals be equally spaced on a line segment of length L , and let traffic be evenly divided among them. If we create a one-level star network (i.e., connect every terminal to the station by a direct and private line), assume that $m \gg 1$ and substitute integrals for sums, we get from (23) that the cost of this one-level system is

$$D = \frac{L}{4} \left[S + \frac{8m}{9T} \right] \quad (24)$$

What would the cost be if we could have used a single line serving a single M/M/1 system?

If we have the same load S , and the average distance a message has to travel is $L/4$ as above, then in order to get the same T from an M/M/1 system our budget will have to be

$$D = \frac{L}{4} \left[S + \frac{1}{T} \right] \quad (25)$$

Defining the quality Q of a system to be the inverse ratio between its cost and the cost of the best possible M/M/1 system, and dividing (25) by (24) we get that the quality of the one-level star system is (for $m \gg 1$)

$$Q = \frac{ST + 1}{ST + 8m/9}$$

Consider now the regular two-level system with P_1 equal groups and P_2 terminals in each group. Assuming that the star network is built at both levels, we get from (23) the following relation between total cost and performance of this two-level system

$$D = \frac{L}{4} S \left(1 + \frac{1}{P_1} \right) + \frac{2L}{9T} \left\{ P_1^{1/2} + P_2^{1/2} \right\}^2 \quad (26)$$

For a given S and T , what should P_1 and P_2 be to minimize D ? Treating P_1 and P_2 as real variables we see that the optimal P_1 and P_2 are related through

$$\frac{P_1}{P_2} = 1 + \frac{9}{8} \frac{ST}{m} \quad (27)$$

The generalization to r levels is immediate. The best r -level regular system must be balanced. That is, $P_1 = P_2 = \dots = P_r = m^{1/r}$ and all individual channels at all levels have the same capacity. The relationship between cost and performance is

$$D = \left(\frac{k}{T} \right)^b r^{b+1} m^{1/r}$$

The best r is easily found to be equal to $\ln(m)/(b+1)$, and when this number of levels is used we get that for all i $P_i = e^{b+1}$ and that

$$D = \left(\frac{k}{T} \right)^b \left(\frac{e}{b+1} \ln(m) \right)^{b+1} \quad (20)$$

Kamoun [15] found similar results when optimizing hierarchical communications networks with other objectives.

If spatial reuse is not perfect and there is some interference between groups we have to modify our formulas slightly. Assume the groups at all but the top level can be colored with q different colors so that no two groups of the same colors at the same level interfere with each other. In an r -level we can now write $T = k \left(\frac{1}{C_1} + \frac{1}{C_2} + \dots + \frac{1}{Cr} \right)$ and $D = P_1 C_1^b + q(P_2 C_2^b + \dots + P_r C_r^b)$. Minimizing D by choosing C_i given T we get $D = \left(\frac{k}{T} \right)^b \left(P_1^{1/(b+1)} + (qP_2)^{1/(b+1)} + \dots + (qP_r)^{1/(b+1)} \right)^{b+1}$. The best P_i satisfy $P_1 = qP_2 = \dots = qP_r$ and using these best values we have

$$D = \left(\frac{k}{T} \right)^b r^{b+1} (mq^{1/(b+1)})^{b+1} \quad (21)$$

We expect q to be a small integer. When m and r grow (21) will give a total cost almost q times greater than that given by (20). But in both cases we see that when using dedicated broadcast channels and the best number of levels the cost of a very bursty system grows like $[\ln(m)]^{b+1}$, and is independent of the geometric dimensionality of the system. The cost of regular hierarchical line networks, given in Theorem 2, depends very much on the dimensionality of the space in which the terminals are distributed. It seems, therefore, that dedicated broadcast channels have a significant advantage over dedicated lines, when building large bursty systems distributed in two or more dimensions.

7. Hierarchical Organization of Non-Bursty Line Systems

So far we have dealt only with extremely bursty systems. Can a hierarchical organization improve the performance of systems that are not bursty?

To answer this question for line networks we have to solve the capacity assignment problem when the traffic is not extremely bursty. This is almost impossible to do explicitly unless the cost of a line is directly proportional to its capacity, which we shall assume in this section. (That is, $b=1$.) Another greatly simplifying assumption we adopt is the *independence assumption* [16]. According to this assumption we analyze the network as if the length of each message is chosen and rechosen *independently*, at each step along its path, from an exponential distribution; and as if arrival of messages at each line is a Poisson process independent of message length. Let C_i , L_i , and S_i be the capacity, length and traffic of the i -th line. The average message delay in getting across the i -th line is then modelled by $T_i = \frac{1}{C_i - S_i}$ and the source-destination delay, averaged over all messages, is $T = \sum S_i T_i / S$. The cost of the i -th line is $D_i = C_i L_i^a$. Minimizing the total cost $D = \sum D_i$ while T is given by choosing C_i , or minimizing T while D is given, we get the following solution for the optimal capacity assignment [16]

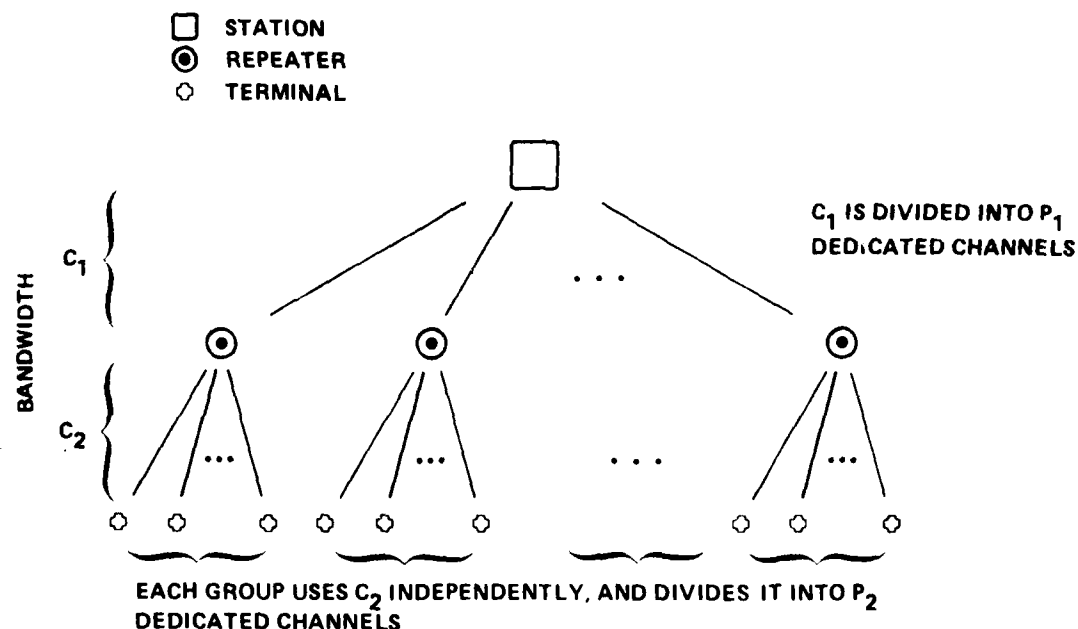


Figure 3. Dedicated Broadcast Channels in a Two-Level Organization.

When the traffic is very bursty the average time spent in this two-level communication system is given by

$$T = k \left(\frac{1}{C_1} + \frac{1}{C_2} \right)$$

where k is a constant depending on the scheme used for splitting a channel into dedicated subchannels. (For Frequency Division Multiple Access $k=1$, for Synchronous Time Division Multiple Access with m subchannels, $k=(m/2+1)/m$). The cost of this two-level system is

$$D = P_1 C_1^b + P_2 C_2^b$$

Our design task is to minimize the necessary budget D , when T and S are given, by choosing C_1 and C_2 , and by choosing P_1 and P_2 subject to $P_1 P_2 = m$.

By symmetry it is obvious that when $m \gg 1$ and two levels are better than one then the best choice is $P_1 = P_2$ and $C_1 = C_2$. That is, the best two-level regular hierarchical broadcast system must be balanced. Using these best values for P_1, P_2, C_1 and C_2 we get

$$D = \left(\frac{k^b}{T^b} \right)^b 2^{b+1} m^{1/2}$$

to the best, i.e., least cost, system. It is quite clear that the concentrator should not be placed in the center of its group but closer to the station. It is quite possible that groups further away from the station should be larger and that messages coming from afar should cross more levels on their way to the station. (This will naturally occur in regular systems too when we note that concentrators will be colocated with some of the terminals, as shown in Fig. 2.)

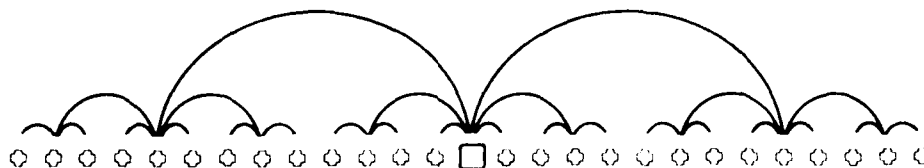


Figure 2. Hierarchical Organization of a One-Dimensional Dedicated Line System.

Some specific heuristics that perturb the regular structure slightly were analyzed in [13], but only a constant improvement was obtained. We suspect that no system will have a cost growing more slowly with m than $m^{1/2}$. (See also discussion at end of section 8.)

6. Dedicated Broadcast Channels

In previous sections we saw that a hierarchical organization can significantly improve the performance of a system based on dedicated lines, especially when the system is bursty. The basic cause for improvement was that instead of having long lines with a small capacity dedicated to each individual terminal we could use short individual lines. The long lines were shared by more traffic, and the capacity invested in them could, therefore, contribute more to improving the performance.

If the communication resource we have is a broadcast channel, whose cost depends on capacity only, it seems that channels used for short distances are just as expensive as those used for long distances. So how can a hierarchical organization help? The crucial fact here is that broadcast capacity can be reused spatially. That is, it can be used independently and at the same time in two or more separate areas. A long range transmission prevents others from using the channel in a large region, and this distance-related 'cost' will be explicitly accounted for in the capacity allocation process.

Let us, once again, create a two-level regular hierarchical system by dividing the m terminals into P_1 groups with P_2 terminals in each. We shall give each group a concentrator, but shall now call it a repeater, this being a more common name when radio networks are discussed [14]. Dedicate a capacity C_1 to every one of the repeater-station communication subchannels. Dedicate a capacity C_2 to every one of the terminal-repeater communication subchannels and assume that these subchannels can be used by every one of the groups to communicate with its repeater, without any interference from other groups. That is, we assume spatial reuse can be done perfectly, without any waste in capacity or degradation in performance. This is a reasonable assumption if, for example, each of the terminals has a directional antenna pointing at its repeater only or if the repeaters are separated by hills, so that every transmission is heard only by the repeater to which it is meant. Fig. 3 shows our model for this two-level system.

Consider the regular hierarchical systems built with a star network at each level. What will be the effect on cost if we change the specification of the allowed delay variance? Consider two extreme cases: In the first, only the average delay is specified. In the second, let us assume that the average delay suffered by messages from any terminal in crossing a given subsystem is the same for all terminals. The comparison between these two alternatives depends on geometric constants, which we have systematically ignored until now since expressing them analytically is usually impossible. To simplify the geometric calculations assume, in this section only, that the region over which terminals are distributed is an n -dimensional sphere, even though a sphere cannot be divided into equal parts similar to itself. Consider first a one-level star network with only the average delay specified. The B -term can be immediately derived from (7). Assuming the number of terminals is large and approximating sums by integrals we get

$$B_A = \frac{n(b+1)}{n(b+1)+a} \left(m S^b L^a \right)^{1/(b+1)}$$

and

$$D_A = \frac{1}{(ST)^b} B_A^{b+1}$$

where the subscript 'A' stands for 'average'.

When a uniform delay is specified D can be written directly, since all channels must have the same capacity, and we get

$$D_U = \frac{m}{T^b} \left[\text{Average of } L_i \right] = \frac{m}{T^b} \frac{n}{a+n} L^a$$

where the subscript 'U' stands for 'uniform'. Forming the ratio we get

$$\frac{D_U}{D_A} = \left[\frac{n(b+1)+a}{n(b+1)} \right]^{b+1} \frac{n}{a+n} \quad (19)$$

Equation (19) was derived by considering one-level systems, but it is valid when comparing r -level systems and when comparing systems with the best r , which is independent of the delay distribution specified. Equation (19) shows, therefore, the additional cost of demanding a uniform delay versus demanding only an average delay.

How large is the ratio given by (19)? It has its largest value when $a=b=n=1$, and is then equal to 9/8. That is, if a system with only the average delay specified is not acceptable, the delay can be made uniform at an additional cost of no more than 12.5 per cent!

5. A Lower Bound?

Theorem 2 shows that by using the heuristic regular hierarchical constructions the cost of very bursty centralized dedicated line systems can be made to grow only slightly faster than $m^{1-a/n}$. (The growth of cost with m can be bounded from above by an exponent of m arbitrarily close to $1-a/n$.) Our regular hierarchical structures have the following properties:

- (1) A concentrator is placed in the middle of each group.
- (2) The terminals are divided and subdivided into equal groups.
- (3) Every message crosses the same number of levels on its way to the station.

These properties were adopted in order to simplify the analysis of regular systems, but they do not lead

The argument of the previous paragraph has the flavor of an existence proof: It shows that by having enough levels the cost can be made to grow as an exponent of m arbitrarily close to $1-a/n$. As m becomes larger, using more and more levels is justified. What is the best number of levels for a large but fixed m ? To answer this question we must consider the constant coefficient multiplying m^k . This constant, which was ignored until now, grows with the number of levels, and therefore tempers the trend towards more and more levels.

The f -function and cost of a system consisting of r levels, each of which is built according to a given design procedure, can be calculated explicitly. Let P_i be the number of terminals per group in the i -th level, starting from the top. Rather than trying to optimize the overall structure directly, note the following: Every two consecutive levels in an optimized r -level system must be optimal as a group of two-level systems. Equation (12) can therefore be rewritten as

$$g^{b+1} P_i^{g-1+a/n} = \left(g - 1 + a/n\right)^{b+1} P_{i+1}^g \quad (15)$$

and (12) can be generalized into

$$\frac{B_i}{B_{i+1}} = 1 - (1-a/n)/g \quad (16)$$

where B_i is the contribution of the i -th level to the B -term. From (15) and (16) we get the following.

Theorem 2: A design procedure for n -dimensional centralized systems whose cost is proportional to m^g where $g > 1-a/n$ can be improved for large m by a multi-level regular organization.

When $1-a/n \neq 0$ the best number of levels is given by

$$r(b+1)/g \ln(g/(g-1+a/n)) = (1-a/n) \ln(m) \quad (17)$$

and the cost of the system, when using this r , is proportional to

$$\left[m^{(1-a/n)/(b+1)} - 1\right]^{b+1} \quad (18)$$

When $1-a/n=0$ the best number of levels is given by

$$r = \frac{g}{b+1} \ln(m)$$

and the cost of the resulting system is proportional to $[\ln(m)]^{b+1}$. In both cases, when the optimal number of levels is used, the number of lines in all groups at all levels is the same, and must therefore be given by $m^{1/r}$.

Proof: See appendix.

When a is smaller the best regular hierarchical system has fewer levels and leads to smaller improvements, since it is harder to save by shortening individual lines. When b is smaller the best system has more levels and leads to larger improvements, since common large capacity lines become more economical.

Example 2: Let the given design procedure be to build a star network. That is, $g=1$. Let a and b be equal to 1. From (17) we see that the optimal number of levels for a two-dimensional system is given in this case by $r = \log_{16} m$, and that we should have 16 lines in every group. The cost of the resulting system is

$$D \approx \frac{L}{T} \left(m^{1/4} - 1\right)^2$$

where we use $\approx$ to denote 'is proportional to'.

$$h = \frac{g^2}{2g - 1 + a/n} \quad (13)$$

When $g > 1 - a/n$ we have $g > h$. That is, when using the best P , as given by (12), we have a two-level structure whose cost grows with m more slowly than the cost of the one-level structure. When $g > 1 - a/n$ and $m \gg 1$ our use of the approximate (10) is consistent, since our best P does satisfy $m \gg P \gg 1$. We can summarize the above discussion of two-level regular hierarchical systems by the following:

Theorem 1: A design procedure whose cost is proportional to m^g where $g > 1 - a/n$ can be improved for large m by applying it separately to each level of a two-level regular structure. The best P (number of groups) is given by (12). The cost of the resulting two-level structure is proportional to m^h , where h is given by (13). When the best P is used, the contribution of the two levels to the delay, to the cost and to the B -term satisfy

$$\frac{T_{top}}{T_{bottom}} = \frac{D_{top}}{D_{bottom}} = \frac{B_{top}}{B_{bottom}} = \frac{g - 1 + a/n}{g} \quad (14)$$

Proof: Substituting (12) in (11) we get $B_{top}/B_{bottom} = (g - 1 + a/n)/g$. The other two equalities are true whenever capacity is optimally allocated, as shown in (8). $\square$

We shall paraphrase (14) by saying that the optimal two-level regular structure is *balanced*. The contribution of both levels to the delay and their share of the budget must be in the proportion given by (14). The right hand side of (14) decreases when g decreases. P also decreases with g , and there will be less groups in the top level. We may say that when g is small most of the system migrates to the bottom level, and that when g is small enough two levels become unnecessary.

Example 1: When the original design procedure consists of building a star network we have $g=1$, and (13) reduces to $h = n/(a+n)$. That is, the cost of the optimal regular two-level star system is proportional to $m^{n/(a+n)}$, while the cost of a one-level system is proportional to m . When $g=1$ (14) reduces to

$$\frac{T_{top}}{T_{bottom}} = \frac{D_{top}}{D_{bottom}} = \frac{a}{n}$$

and we get that the two levels must be balanced in a way that depends on the dimensionality of the system and on the economy of scale of long lines, but is independent of the possible economy of scale involving capacity. $\square$

If two levels are good, will more levels be better? Equation (13) already contains the answer: Decomposing a given system into two levels and applying the original design procedure to each can be considered as a new design procedure. Applying this new procedure to two levels is equivalent to applying the original procedure to four levels. When $g > 1 - a/n$ it follows from (13) that $h > 1 - a/n$ and therefore four levels will be better than two when m is large enough. In general, let g_0 be the power of m characterizing the resulting cost and f -function when the given design procedure is applied to 2 levels. Equation (13) can be rewritten as

$$g_{i+1} = \frac{g_i^2}{g_i - 1 + a/n}$$

where g_0 is the power of m characterizing the direct application of the given design procedure to one level. It is easy to see that when $g > 1 - a/n$ the sequence $\{g_i\}$ is monotonically decreasing and converges to $1 - a/n$.

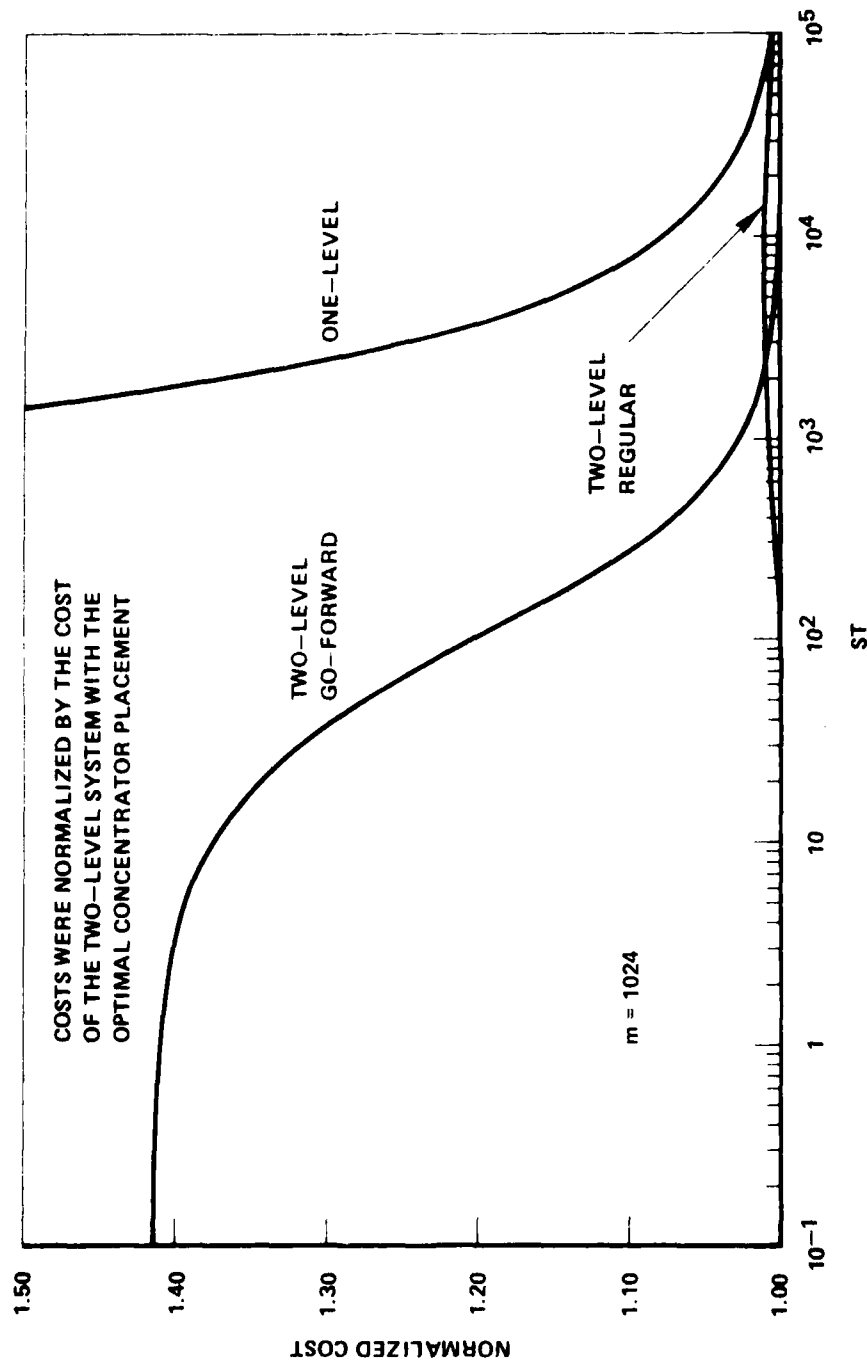


Figure 8. The Improvement Achieved by Optimizing Concentrator Placement.

Consider a network in which each terminal is connected to its nearest neighbor on each side, as shown in Fig. 9.

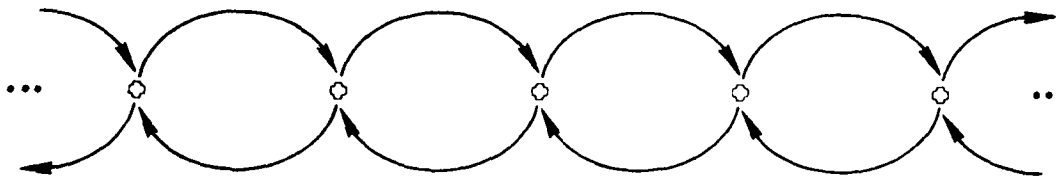


Figure 9. The One-Level One-Dimensional Network.

Let C be the capacity given to each line. We shall call this network the one-level network. Every message goes, on the average, through N/l lines on its way from source to destination, and the traffic in each line is $NS_u/2$. Hence, the average message delay is given by

$$T = \frac{N/l}{C - NS_u/2} \quad (31)$$

Assume that the cost per unit length of a half-duplex line is equal to its capacity, i.e., $a=b=1$. To calculate the budget per unit length necessary for satisfying a given T and S_u via the one-level system, we solve (31) for C as a function of N , S_u and T , and then multiply by two, since every unit interval has exactly one line carrying traffic in each direction. The result is

$$D_u = NS_u + 2M/T \quad (32)$$

where $M=N/l$ is the number of terminals contained in the average path.

It seems that NS_u , the traffic coming out of a portion of the network whose length is equal to the average distance travelled, is a natural traffic measure in a one-dimensional network. (After all, S_u has the dimensions of traffic per length, and what other natural length except N do we have to multiply S_u with in order to get something with the dimension of traffic?) The natural dimensionless parameter we shall use to characterize the traffic is $NS_u T$. When it is small we shall call the traffic bursty, and when it is large we shall call the traffic steady.

Let us double the number of terminals per unit length, while keeping the traffic per unit length, and the average distance travelled by messages constant. Each terminal will now generate half the traffic a terminal generated in the original system. The new network has the same N and S_u , but M became twice as large. M plays in (32) the same role that m played when we discussed centralized systems. It is a natural measure for the network being distributed, and characterized the extra expense incurred because terminals are not all at one place. We conclude from (32) that the fact that the network is distributed poses no problem when the traffic is steady. (When $NS_u T \gg 1$ the second term in (32), which is the only one that depends on M , is negligible compared with the first.) But when the traffic is bursty, the system cost is essentially proportional to M , i.e., the cost is then strongly dependent on how distributed the system is.

Can hierarchical organizations help networks? Can we use concepts introduced previously for centralized systems to characterize good hierarchical networks?

9. Hierarchical Line Networks

Consider now n -dimensional networks in which the cost of a line is, once again, equal to the product of the a -th power of its length times the b -th power of its capacity. Let N be the average source-destination distance to be travelled by messages, and assume the size of the networks is much larger than N , so that edge effects can be neglected. Let M be the number of terminals in an n -dimensional cube of size N . The volume occupied by every terminal has therefore a typical linear size equal to $N/M^{1/n}$.

Let us form a one-level network by connecting every terminal to a small number of near neighbors. The typical line length is $N/M^{1/n}$, and every message typically goes through $M^{1/n}$ lines. The cost per unit volume is therefore given by

$$D_u \approx \frac{M}{N^n} \left(\frac{N}{M^{1/n}} \right)^a \left(\frac{M^{1/n}}{T} \right)^b = \frac{N^{a-n}}{T^b} M^{(n+b-a)/n}$$

To build a hierarchical system we shall introduce *stations*, connect every station to a few of its near neighbors, and assume messages are routed thus: Every message will go from its source terminal to the nearest station, from it to the station nearest its destination using the inter-station lines, and from that final station to its destination. Let L be the length of the typical inter-station line, and let $1/L^n$ be the density of stations.

When networks are very distributed (i.e., $M \gg 1$) a good placement of stations will usually satisfy $N \gg L \gg N/M^{1/n}$. We shall call the inequality $N \gg L$ the assumption of *long distance travel* and consistently use two of its implications: The portion of traffic that can reach its destination without getting to any station is negligible, and the average line of sight distance travelled by a message from the station near its source to the one near its destination is approximated well by N . The assumption of long distance travel allows us therefore to ignore the distribution of distance travelled. Considering this distribution is of no importance when optimizing a multi-level structure with $M \gg 1$ [3].

If we assume that every terminal is connected to its station by a direct line we get a two-level system. Using the assumption of long distance travel we can calculate its cost thus: Let $1/T_1$ and $1/T_2$ be the typical capacity of lines in the inter-station (top) level and the terminal-station (bottom) level respectively. A typical message takes 2 hops on lines in the bottom level and N/L hops in the top level (ignoring a small geometric constant.) The average time a message spends in getting from source to destination is therefore

$$T \approx \frac{N}{L} T_1 + 2 T_2$$

Let there be p terminals per station. The typical length of lines in both levels is L , and the cost per unit volume is therefore

$$D_u \approx \frac{M}{N^n} \left(\frac{1}{p} \frac{L^a}{T_1^b} + \frac{L^a}{T_2^b} \right)$$

minimizing D_u by choosing T_1 and T_2 given T we get, in analogy to (7),

$$D_u \approx \frac{M}{N^n} \frac{L^a}{T^b} B^{b+1}$$

where $B = (N/L)^{b/(b+1)} / p^{1/(b+1)} + 2^{b/(b+1)}$. The p that will minimize D_u must satisfy

$$\left(\frac{n+b-a}{a} \right)^{b+1} (M/p)^{b-n} \approx 2^b p \quad (33)$$

When this best p is used we have, independently of the geometric constants neglected when writing (33), that

$$\frac{D_{top}}{D_{bottom}} = \frac{B_{top}}{B_{bottom}} = \frac{a}{n+b-a} \quad (34)$$

$$D_u \approx \frac{N^{a-n}}{T^b} M^{(b+n-a)(b+n)}$$

Equation (34) shows, once again, that the best two-level system is balanced, but the optimal investment ratio for networks, given in (34), is different from the optimal investment ratio in centralized systems, given in (14).

r -level networks, with $r-1$ levels in the terminal-station part, can be solved by applying (33) and (34) to the top two levels, and by applying (15) and (16) to any other two consecutive levels. But the network with the best number of levels can be more simply characterized by applying Theorem 2 to the terminal-station part, i.e., by assuming that every one of the centralized systems connecting terminals to their station has the best number of levels. Assume that the inter-station distance is L and that the number of terminals per station is p . When $a \neq n$ we get from (18) that the cost per unit volume of the terminal-station levels is

$$\frac{1}{L^n} \frac{L^a}{T^b} \left(p^{(1-a-n)/(b+1)} - 1 \right)^{b+1} \quad (35)$$

Using $nN^n = ML^n$ to express L in terms of p we see that when $p \gg 1$ this cost is a slowly growing function of p , proportional to $p^{a/n-1} \left(p^{(1-a-n)/(b+1)} - 1 \right)^{b+1}$. The top-level cost is, when $p \gg 1$, a slowly decreasing function of p , and the best p is therefore of a magnitude similar to M . When the traffic is very bursty and $M \gg 1$ the cost per unit volume of a network with the best number of levels can therefore be roughly given by

$$D_u \approx \frac{N^{a-n}}{T^b} M^{1-a-n} \quad (36)$$

Continuing the discussion of a possible lower bound for the cost of line systems started in section 5 we can say the following: If centralized systems existed whose cost grew more slowly with p than p^{1-a-n} then instead of (35) we would have that the cost of the terminal-station levels is a *decreasing* function of p . The overall network cost would then be a decreasing function of p and of L and the best L will satisfy $L \gg N$. While not impossible, it is very strange that the best network will force a message to go to a station that is much further away from its source than is its average destination.

In analogy to (36) one can see [3] that the cost of very bursty broadcast networks and of one-dimensional line networks is proportional to $[\log(M)]^{b+1}$, and broadcast channels are once again superior to lines for a bursty system distributed in more than one dimension.

10. The Geometry of Networks

In deriving (36) we neglected various geometric constants, since we wanted to show in the simplest possible form how the cost of very bursty networks depend on system parameters. (While (36) does not contain S_u , it is valid only when $N^n S_u T \ll 1$.) How will the geometry of the top level influence the cost of networks? We shall treat only the case of two-dimensional line networks.

It is well known [17] that there are exactly three regular tessalations of the plane: i.e., three ways to cover the plane with identical regular polygons. If we place a station in the middle of each tile and connect it to its nearest neighbors we get the three networks shown in Fig. 10. We shall call them the square, triangular and hexagonal network, where the name applies to the regular polygons created by the lines in the network. Note that we do not draw the *tiles* (the regions around each station), but the dual graph showing the communication lines between adjacent stations. For example, tessalating the plane by hexagonal tiles produces a beehive-like structure which leads to our triangular networks.

Is there a common basis for comparing these three tessalations? For a preliminary comparison, let us assume that all traffic originates at the stations, and is destined to many points in the plane, not necessarily to other stations in the network. Every message will use the given network to arrive at the node closest to its destination. We shall not consider how the final node delivers each message to its exact destination at this time. Let us also assume that the distribution of traffic coming out of a node has a radial symmetry, and that the average line of sight distance from the source node to the destination node is N . The average distance actually travelled by a message will be larger, because there may not be a line directly to the neighborhood of its destination. Assuming the average distance travelled is much larger than the inter-node distance we can say that the distance actually travelled is δN , where δ is a characteristic constant for each of the possible networks.

In the square network we have $\delta = \frac{1}{2\pi} \int_0^{2\pi} (|\cos\theta| + |\sin\theta|) d\theta = \frac{4}{\pi} = 1.27$. A similar simple calculation gives that in a triangular network $\delta = 1.10$. For the hexagonal network we used a computer program to find that δ is approximately equal to 1.30.

Let S_u and D_u denote the total traffic and budget per unit area. Let A be the area per node. Each node will generate new messages at a rate of AS_u . If L is the internode distance then the number of hops taken by a message, on the average, is $\delta N/L$. Therefore the total traffic passing through each node will be $AS_u \delta N/L$ messages per second. Let ϵ be the number of nearest neighbors each node has, which is also the number of (half-duplex) lines per node. The total traffic per line must therefore be $AS_u \delta N/L\epsilon$. If T is the required total average delay, the delay suffered when crossing a given line must be $TL/\delta N$, and the capacity necessary for each line is

$$C = AS_u \frac{\delta N}{L\epsilon} + \frac{\delta N}{LT} \quad (37)$$

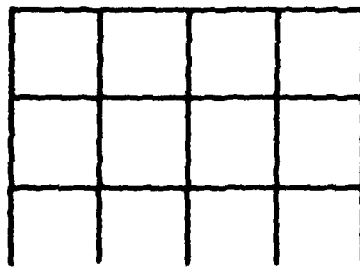
Let us first assume that $a=b=1$. The total cost per node is then found by multiplying (37) by L , the length of every one of the lines, and by ϵ , the number of lines per node. If we divide by A , we find the cost per unit area to be:

$$D_u = \delta NS_u + \frac{\epsilon \delta N}{TA} \quad (38)$$

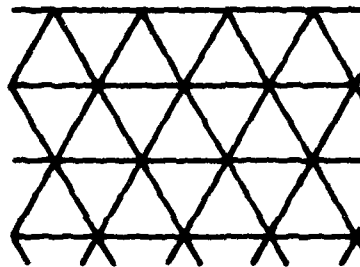
Let M be the number of nodes in a square whose sides are equal to N . The area per node is then N^2/M . Substituting this for A in (38) we get

$$D_u = \delta NS_u + \frac{\epsilon \delta M}{N} \frac{1}{T} \quad (39)$$

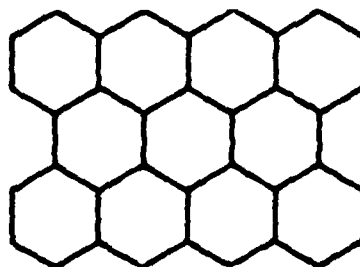
In a two dimensional network, the natural traffic measure is $N^2 S_u$, and the burstiness measure is $N^2 S_u T$. When the traffic is very steady only the first term in (39) is significant. The best network will then be the one with the smallest δ , i.e., since the triangular network imposes the least extra distance on messages, it is the best of the three for steady traffic. When the traffic becomes very bursty ($N^2 S_u T \ll 1$) only the second term in (39) is significant, and the hexagonal network is the best because it has the smallest $\epsilon \delta$.



SQUARE



TRIANGULAR



HEXAGONAL

Figure 10. The Three Regular Tessellations of the Plane.

Fig. 11 shows the cost of the three networks, normalized by the cost of a hypothetical network in which $\delta=\epsilon=1$. As expected, the triangular network and the hexagonal network are best when the traffic is, respectively, very steady and very bursty. It is somewhat surprising, though, that the square network is never the cheapest of the three.

For general a and b , i.e., not necessary equal to 1, we get from (37)

$$D_u = \frac{1}{A} \epsilon L^a \left(AS_u \delta N / L \epsilon + \delta N / TL \right)^b$$

L does not in general disappear from the cost formula, but we can write $A = \eta L^2$, where η is a constant, depending on the geometry of the network, and given in Table 1. When comparing the three regular networks we shall assume that A and the density of terminals are common to all three. We then find the following: When the traffic is very bursty the best network is the one having the smallest $\epsilon \eta^{(b-a)/2} \delta^b$. When the traffic is very steady, the best network is the one having the smallest $\epsilon^{1-b} \eta^{(b-a)/2} \delta^b$.

It is quite intuitive that as b grows smaller the advantage of the hexagonal network grows, since it concentrates its traffic on fewer high capacity lines that are becoming relatively cheaper. As a grows smaller the advantage of the hexagonal network decreases, since its line channels are shorter. Using the numeric values for δ , ϵ and η we find that of the three regular networks, the hexagonal is always (i.e., independently of a and b) the best when the traffic is very bursty. When the traffic is very steady the hexagonal network is better when $b \leq 0.65 + 0.19a$, otherwise the triangular is better.

There is, of course, no reason to limit our consideration to the three networks in which all nodes are equivalent and in which lines connect only nearest neighbors. When the traffic is steady, we can connect every node to more of its neighbors, in order to lessen the distance messages have to travel. However, since the triangular network already has $\delta=1.10$, the most we can gain by introducing more and more lines is 10%. When the traffic is bursty there is room for a lot of improvement, and that is where hierarchical structures become interesting.

Newell [18] gives a general discussion of networks with an economy of scale in their cost. He points out that even if the node placement and the traffic requirements are symmetric, the best network will in general *not* have the same symmetry. For example, the two-dimensional square network with a large M and a bursty traffic can be improved by deleting every other vertical line. The resulting structure, shown in Fig. 12, forces some messages to go an extra distance, until they can find a vertical line. But as a result only half as many vertical lines are necessary, and when the traffic is bursty this will more than compensate for the extra distance travelled.

In our model there can be three independent sources for an economy of scale: when $b < 1$ large capacity lines are relatively cheaper, when $a < 1$ long lines are relatively cheaper, and when the traffic is bursty sharing unused resources leads to significant economies. What is the best network structure, as a function of a , b and burstiness? Newell, in the same paper quoted above [18], points out that there are no efficient algorithms for solving large minimization problems when the cost functions are concave, i.e., when there is an economy of scale. Symmetry cannot be used to reduce the complexity of the problem, because the best solution will not necessarily reflect the symmetry of the traffic requirements. We shall not, therefore, try to find the best network. Can any conclusions be drawn by considering the geometry of our heuristically constructed hierarchical structures? In the previous section we ignored the geometric constants, but let us now bring them into the treatment of two-level networks, when $a=b=1$ and when the traffic is not necessarily bursty.

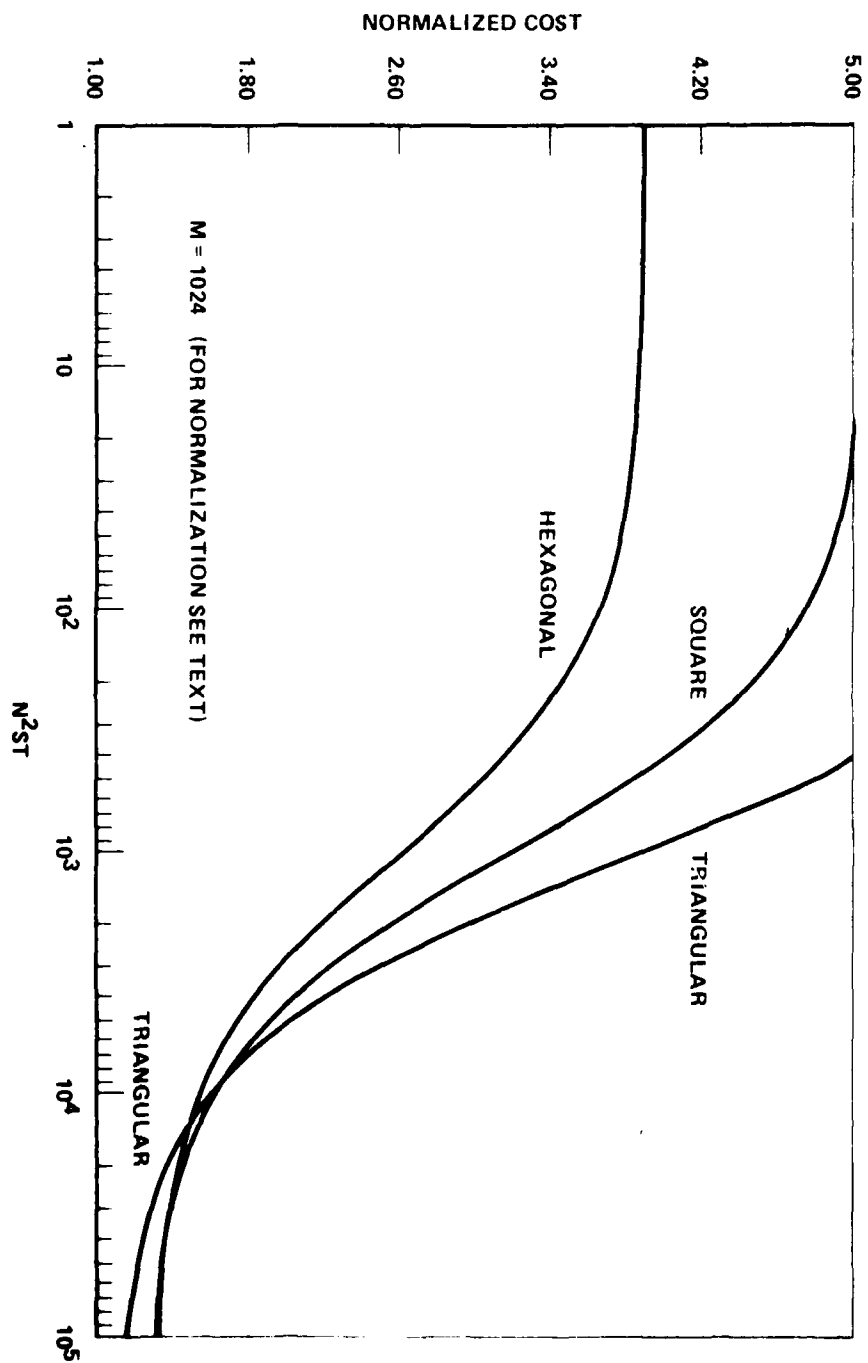


Figure 11. Cost of One-Level Regular Networks.

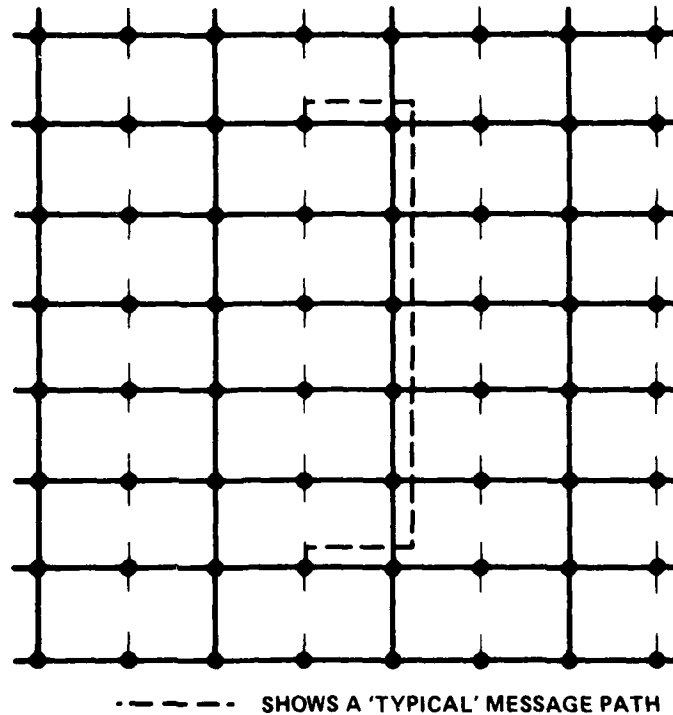


Figure 12. Improving a Square Network for Bursty Traffic.

Let D_1 be the cost per unit area of the top level (the station-station level), and let T_1 be the average time each message spends in the top level. (38) applies directly to the top level. Let L be the distance between nearest stations, and let A be the area per station, where $A = \eta L^2$. By definition, N is the average line-of-sight terminal-to-terminal distance a message has to travel. When $N \gg L$, N is also the average station-to-station distance a message has to travel. Therefore, from (38), the cost per area of the top level is given by

$$D_1 = \delta NS_u + \frac{\epsilon \delta N}{\eta L^2} \frac{1}{T_1} \quad (40)$$

In order to calculate the cost of the bottom level, we must find the average terminal-station distance. If the area assigned to a station was a circle of diameter L this average distance would have been $L/3$. In practical networks with inter-station distance L the average terminal-station distance must be larger, and we shall write it as $\zeta L/3$, where ζ is a constant to be determined. The average square root of the terminal-station distance will similarly be written as $\xi(4/5)\sqrt{L}/2$, where ξ is a constant.

A summary of the numerical coefficients characterizing the networks built with the three regular tessalations at the top level is given in Table 1. Also included in the table is the hypothetical, but impossible, 'best' network, which we use for normalizing the cost in our figures.

Table 1
Coefficients Characterizing the Geometry of Two-Dimensional Networks

		square	triangular	hexagonal	'best'
actual distance/line of sight	δ	1.27	1.10	1.30	1
lines per node	ϵ	4	6	3	1
area per node	η	1	$\sqrt{3}/2$	$3\sqrt{3}/4$	$3\sqrt{3}/4$
terminal-station distance	ζ	1.148	1.05	1.38	1
terminal-station $\sqrt{\text{distance}}$	ξ	1.070	1.026	1.168	1

Let D_2 be the cost per unit area of the terminal-station level, and let T_2 be the average time each message spends in this level. Since every message goes through this bottom level twice, once at each end of its path, and since each terminal has two half-duplex lines, for sending to and receiving from the station, respectively, we see that

$$D_2 = 2S_u \zeta \frac{L}{3} + \frac{32M}{25N^2} L \xi^2 \frac{1}{T_2} \quad (41)$$

where, as before, M/N^2 is simply our way of writing the terminal density. For a given L the total cost of the two-level network can be obtained from (40) and (41) when minimizing $D_1 + D_2$ by choosing T_1 and T_2 subject to $T = T_1 + T_2$. Let x be the ratio between L and N , that is, x is the interstation distance measured in the natural distance unit of our networks. The total cost is then

$$D_u = NS_u(\delta + 2\zeta x/3) + \frac{1}{T} \left[\frac{1}{x} \sqrt{\epsilon \delta / (\eta N)} + \frac{4}{5} \xi \sqrt{2Mx/N} \right]^2 \quad (42)$$

Equation (42) gives the total cost of a two-level two-dimensional networks as a function of x , the ratio between the interstation distance and the average distance travelled by a message. Which x will minimize D_u ? This best x is easily found numerically, and Fig. 13 shows the cost of two-level systems, in which the top level was a square, triangular or hexagonal network. The cost of these networks, where the best x was chosen for each as a function of $N^2 S_u T$, was normalized by the cost of the hypothetical 'best' network defined by Table 1, with its best x as a function of $N^2 S_u T$.

Once again, we see, that the square network is never the best. When the traffic is bursty the hexagonal network is best, and when the traffic is steady the triangular network takes over. Comparing Figs. 11 and 13 we see that in the two-level system the triangular network becomes better than the hexagonal one at a smaller value of $N^2 S_u T$ than in the one-level system. This is because we simply ignored the question of how messages arrived at the stations in our treatment of one-level networks. In our model for two-level networks we explicitly took into account the terminal-station distance. If we compare our three networks with the same area per node we see that the triangular network has the smallest average terminal-station distance, and the hexagonal network has the largest average distance. This distance is irrelevant when the traffic is bursty, but gradually becomes important as the traffic becomes steady, and is the reason for the earlier superiority of two-level triangular over hexagonal networks.

Figs. 11 and 13 are both drawn for $M=1024$. If we consider a different M the one-level curves of Fig. 11 will simply be shifted along the $N^2 S_u T$ axis, while retaining their shape. The shape of the curves describing the two-level networks is not invariant when M changes, but the general characteristics were checked for $M=16, 256, 1024, 4096$ and 16384 , and they are the same: triangular two-level networks are good for steady traffic, hexagonal networks are good for bursty traffic, and the square networks are never the best of the three.

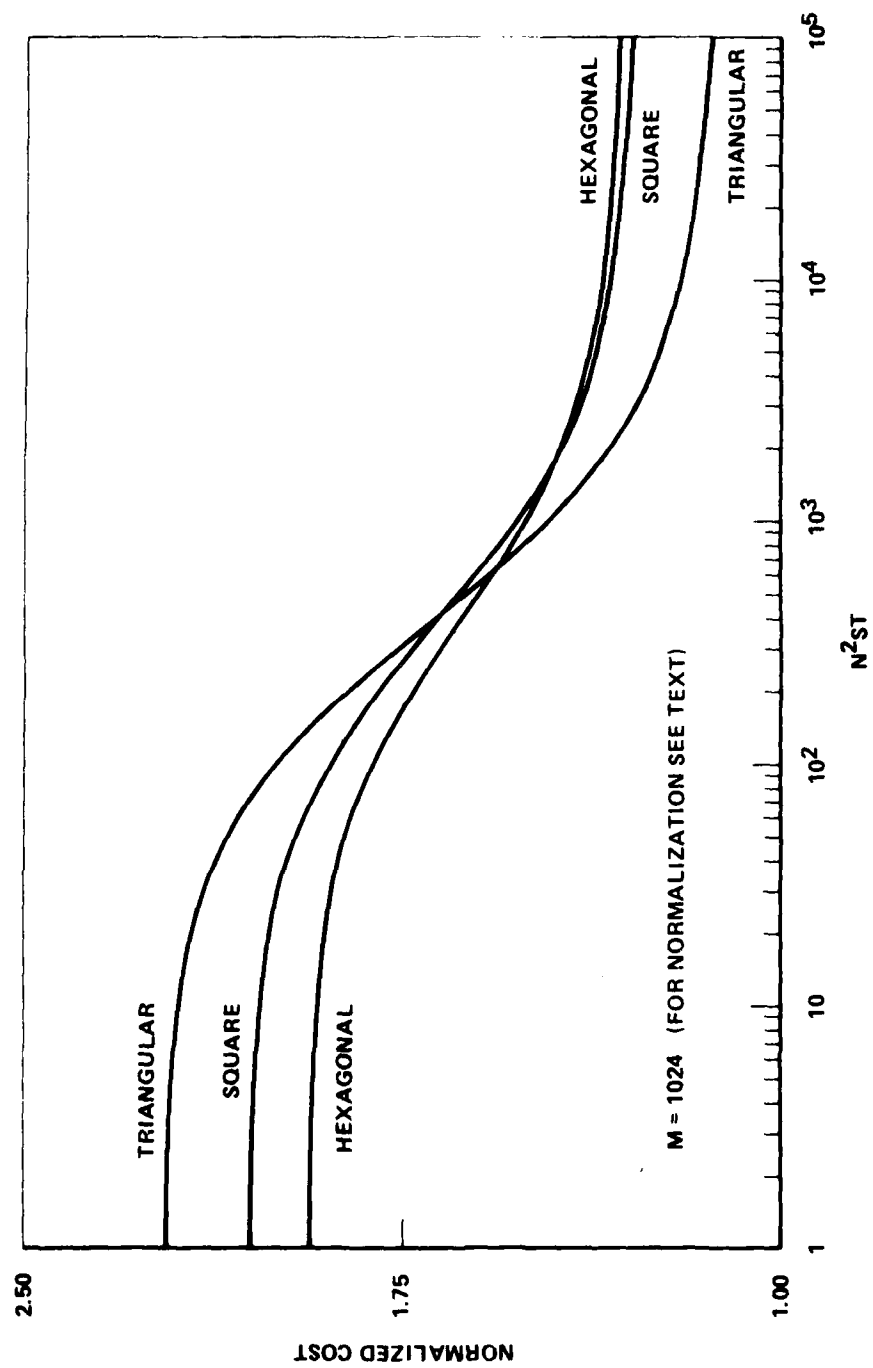


Figure 13. Cost of Two-Level Networks.

When the traffic is bursty and M is large, more than two levels will be even better. What will be the good geometry? In the previous section, while deriving (36), we saw that when the traffic is very bursty and when $M \gg 1$ the network cost is dominated by the terminal-station part. It immediately follows that the best network will be the triangular, which has the smallest terminal-station average distance when the area per station is given. Combining this conclusion with the previous discussion of Figs. 11 and 13 we are tempted to conjecture that whenever the best number of levels, as a function of burstiness, is used, the top level should have the triangular geometry. The top level will either be steady enough, or else it will be just one of many levels, and the cost of all but the top level will make our triangular network (with its hexagonal tiles) the best. For the same reason it is natural to assume that the top level will always reflect the translational and rotational symmetry of the traffic requirements, and that we shall never have to use networks like that of Fig. 12 in the top level.

11. Conclusions

We have assumed that the traffic level and the necessary performance are specified, and that the goal is to fulfill these requirements with the least cost. The quality of a given organization is defined to be the inverse of the cost of a given organization, suitably normalized. Burstiness is defined and serves as a natural dimensionless number to characterize the requirements. We also assume that space is homogeneous and isotropic: terminal density and traffic requirements are the same everywhere. This often leads to results that depend only on the average distance travelled by messages, and not on the distribution of distances travelled. The validity of our results in the case of irregularity either in spatial distribution or in traffic requirements was not investigated. The cost of communication resources was modelled by simple power laws.

When the traffic is steady, the quality of simple one-level dedicated-channel systems is reasonably good, since all channels will be well utilized. When the traffic is bursty, channels are hardly utilized, and a significant gain can be achieved by sharing, even if the technology has no inherent economies of scale.

To make sharing of dedicated channels possible, we introduce *regular* hierarchical structures. (For a treatment of hierarchical organization mixing ALOHA and dedicated channel see [19].) Our regular structures are obtained by dividing the terminal population into equal groups, and placing a concentrator in the center of each. Regular multi-level hierarchical structures can improve the performance of bursty systems significantly. The optimal structure is characterized by a balance principle, that gives the ratio of investment in any two consecutive levels. Another characteristic of the optimal regular hierarchical structures is that channels are organized in small groups of equal sizes.

In line systems the improvement is obtained by shortening individual lines and from sharing long high-capacity lines. The performance of regular line structures is therefore strongly dependent on the dimensionality of the system. It is harder to improve two and three-dimensional line systems by our regular structures since the typical line length decreases more slowly with the number of groups when the terminals are distributed in more dimensions. The question of the performance of the best possible line structure is raised but left open. We conjecture that the dependence of the cost of regular structures on dimensionality will not be significantly improved by any scheme.

The improvement of broadcast systems follows from spatial reuse: i.e., different groups of terminals can communicate with their concentrators by short range transmissions at the same time, thereby sharing bandwidth. The performance of regular broadcast systems is independent of dimensionality, and very similar to that of the one-dimensional line systems. For systems in two or more dimensions which are very distributed and bursty, dedicated broadcast channels are therefore better than line channels.

The problem of very bursty distributed networks with dedicated channels reduces almost entirely to the centralized system problem, since the 'network' part at the top level is only one of very many levels. Tessalating the plane with hexagonal tiles leads to the best network with both technologies, but for different reasons. Of all regular shapes tessalating the plane, the hexagon has the smallest average distance to its 'center', and this makes it superior for line networks. Tessalating with hexagons is good for broadcast networks using omnidirectional antennas because it results in the least interaction between neighboring tiles, and makes the most sharing possible [3].

The best geometry for a network with a given number of levels changes with burstiness, but it seems that, for line networks, when the best number of levels is used, as a function of burstiness, tessalating the plane with hexagonal tiles (and forming a triangular network of communication lines) is usually the best.

Appendix

To simplify our formulas here let us rewrite (15) and (16) as

$$P_i = P_{i+1} \quad (A1)$$

$$\frac{B_i}{B_{i+1}} = \frac{1}{x} \quad (A2)$$

where

$$x = \frac{g}{g-1+a/n}$$

$$s = 1/x$$

$$t = x^{(b+1)/g}$$

Using $\prod P_i = m$ and $\sum B_i = b$ we can solve (A1) and (A2) for P_i and B_i in terms of B, r, t, s, x and m .

When $a \neq n$ we get

$$P_i = t^{1/(1-s)} \left(m^{1-s} t^{-r} \right)^{\frac{s-1}{1-s}} \quad (A3)$$

$$B_i = x^{i-1} \frac{1-x}{1-x^r} B \quad (A4)$$

Ignoring geometric constants, we also know that the following must be true

$$B_i \approx \left(L^a S^b P_i^g \right)^{\frac{1}{b+1}} \quad (A5)$$

Using (A3) and (A4) in (A5) we can get B as a function of m, r and the constants t, s and x . Isolating the dependence on r we get that B is proportional to

$$(x^r - 1) \left(m^{1-s} t^{-r} \right)^{\frac{1}{1-s} - \frac{g}{b+1}} \quad (A6)$$

Differentiating we find that B is minimized, as a function of r , when

$$m^{1-s} = x^{r(b+1)/g} \quad (A7)$$

Substituting (A7) in (A6) we get that B is proportional to $(x^r - 1)$ and is therefore proportional to $m^{(1-s)g/(b+1)} - 1$. Since the cost is proportional to B^{b+1} it follows that when the best r is used the system cost is proportional to

$$\left(m^{\frac{1-a}{b+1}} - 1 \right)^{b+1} \quad (\text{A8})$$

Substituting (A7) in (A3) we also see that when the best r is used the P_i do not depend on u , and they must therefore satisfy $P_i = m^{1/r}$. While the best number of levels will depend on g , i.e., on the quality of the design procedure applied to each level, (A8) shows that the system cost, when the best number of levels is used, is independent of g . For larger m we can also approximate (A8) by $m^{1-a/n}$ and see that the growth with m of the best regular hierarchical system depends only on the geometric dimensionality and on the length dependence of line cost, and hardly depends on the capacity dependence of line cost.

When $a=n$ (A3)-(A7) are not valid since $x=s=r=1$. But the solution is actually simpler. In this case we get from (A1) that for every r , the best r -level system should have $P_1=P_2=\dots=P_r=m^{1/r}$, and from (A2) we get that $B_1=B_2=\dots=B_r=B/r$. Substituting in (A5) and ignoring the geometric constants we get

$$B/r \approx \left(L^a S^b m^{s/r} \right)^{1/(b+1)}$$

Isolating the r -dependence of B , it is easy to see that the best r must satisfy $(b+1)r = g \ln(m)$, and that the system cost when the best number of levels is used is proportional to $[\ln(m)]^{b+1}$.

References

- [1] Kleinrock, L., *Queueing Systems, Vol. I: Theory*, Wiley-Interscience, New York, 1975.
- [2] Kleinrock, L., "Resource Allocation in Computer Systems and Computer-Communication Networks," *Information Processing 1974, Proceedings of IFIP Congress '74*, Stockholm, August 1974, North Holland, Amsterdam 1974, pp. 11-18.
- [3] Akavia, G.Y., "Hierarchical Organization of Distributed Packet Switching Communication Systems," Ph.D. Dissertation, Computer Science Department, University of California, Los Angeles, California, 1978.
- [4] Ferguson, M.J., "A Bound and Approximation of Delay Distribution for Fixed-Length Packets in an Unslotted ALOHA Channel and a Comparison with Time Division Multiplexing (TDM)," *IEEE Transactions on Communications*, Vol. COM-25, pp. 136-139, January 1977.
- [5] Lam, S., "A New Measure for Characterizing Data Traffic," *IEEE Transactions on Communications*, Vol. COM-26, pp. 137-140, January 1978.
- [6] Martin, J., *Systems Analysis for Data Transmission*, Prentice-Hall, Englewood Cliffs, New Jersey, 1972.
- [7] Roberts, L.G., "Dynamic Allocation of Satellite Capacity through Packet Reservation," *AFIPS Conference Proceedings*, 1973 National Computer Conference, Vol. 42, pp. 711-716.
- [8] Binder, R., "A Dynamic Packet-Switching System for Satellite Broadcast Channels," *Proceedings of the IEEE International Conference on Communications*, San Francisco, California, pp. 41-1 to 41-5, June 16-18, 1975.

$$P_1 = \sqrt{m} + \frac{S}{C-S}, P_2 = \frac{m}{P_1}$$

and the delay is then given by

$$T = \frac{S}{(C-S)^2} + \frac{2\sqrt{m}}{C-S} \quad (9)$$

Comparing (9) with (7) we see that we gain a factor of at least 2 by going from a two-level channel-splitting scheme to a two-level channel-sharing scheme. Comparing (9) to the delay in the one-level scheme, given by $T = \frac{m}{C-S}$, we see that the two-level channel-sharing scheme is better than the one-level scheme as long as P_1 is smaller than m .

6. Mixing with a General Random-Access Scheme

The channel-splitting curve in Fig. 5 shows the power of the two-level mixed-mode idea even in its simplest form: by introducing intermediate repeaters and choosing their number we gain a significant improvement over both the one-level ALOHA and the one-level FDMA. By choosing the number of repeaters, we can make sure that the dedicated channels are not underutilized and that we do not have ALOHA systems that are too heavily loaded.

Retracing our steps so far, we can see two ideas that improve the mixed-mode organization even more: If top-level transmission can be perfectly captured in the presence of bottom-level transmission then both levels should share the channel, and we get the 'full interference' case. If the interaction between levels is minimal then the performance is even better, and the 'no interference' model is then appropriate.

Fig. 6 repeats some of the curves of previous sections and also includes the two-level ALOHA scheme of [4]. We see that two-level ALOHA offers little improvement over the two-level mixed-mode scheme, even though the two-level ALOHA was modelled with the best possible assumptions regarding the interaction between levels. It can be shown that the three-level ALOHA offers even less improvement [4]. We thus reach the conclusion that if you were to design a system for a given ST where neither one-level ALOHA nor FDMA perform well, you should almost always use a two-level *mixed-mode* system, and only rarely (i.e., for a small range of ST around 1) should you use two-level ALOHA. Intuitively, a message should (almost) never have to face contention systems *twice* on its way to its destination: if contending once is not enough to reach the destination, the rest of the way should consist of dedicated paths.

The dedicated-channel scheme can be improved by a multi-level organization that uses dedicated channels at all levels [3]. Even with the best number of levels, the cost of a multi-level dedicated-channel scheme grows with the number of terminals. The mixed-mode scheme presented here already assumes the population of terminals is 'infinite', and its cost is independent of the number of terminals. A hierarchical organization mixing modes is therefore better, when the number of terminals is large, than a hierarchical organization using dedicated channels only.

Will the analysis presented so far be useful if we have the option of using Carrier Sense Multiple Access (CSMA) or any other random access that is better than ALOHA?

We shall describe a general random access scheme by

$$T = \frac{1}{C - S/U}$$

where U is its maximum utilization. If U is greater than $1/e$ the random access scheme will be better than ALOHA, and the region (in the Quality versus ST plane) left infeasible will become smaller. But

organization and made unavailable to the bottom level was left idle. In the channel-sharing system everything that is not actually used by the top level is available to the bottom level. The new bottom level has therefore more capacity, and the delay in it will be smaller. We thus have that the total delay in the channel-sharing system is smaller than the delay in the channel-splitting system.

Are the assumptions used in proving theorem 1 reasonable? The first is simply the assumption of 'transparent bottom', introduced and justified earlier. When the two levels are synchronized, the total capacity available to the bottom level will be reduced exactly by the amount of activity in the top level. But the assumption that the delay will simply depend on this reduced capacity ignores the details of the occurrences following a transmission failure (for example, the retransmission policy and its influence on delay). The second assumption is thus more a device to approximate and simplify the behavior of real systems than a direct description of them. It is a natural extension of another device we have used consistently: the assumption that the total offered traffic in an ALOHA system is a Poisson process.

The simple model of the influence of the top level on the bottom level, which is assumed in theorem 1, has been used systematically in earlier sections of this paper. As a different example of the benefit of sharing, let us see the improvement possible when dedicated broadcast channels are used.

Assume we have m terminals and form a two-level system by splitting them into P_1 groups with P_2 terminals each. If dedicated channels are used at both levels and there is no interference between lower-level groups we have

$$T = \frac{P_1}{C_1 - S} + \frac{P_2}{C_2 - S/P_1}$$

If the total communication capacity we have is C , the task of designing the best system can be formulated thus: Minimize T when S is given by choosing P_1 and P_2 subject to $P_1 P_2 = m$, and by choosing C_1 and C_2 subject to $C_1 + C_2 = C$. The constrained minimum is achieved when $C_1 = (C+S)/2$, $C_2 = (C-S)/2$ and

$$\frac{P_1}{P_2} = \left[1 + \frac{1}{\sqrt{m}} \cdot \frac{2S}{C-S} \right]^2$$

and the resulting minimum T for a two-level dedicated channel scheme is given by

$$T = \frac{4S}{(C-S)^2} + \frac{4\sqrt{m}}{C-S} \quad (7)$$

What will be the system performance if the channel is shared between levels? To analyze this case we shall assume that the bottom level is transparent and can detect its failures immediately. The lower level uses the empty slots left by the upper level in a round-robin fashion. The total delay for a system of P_1 groups of P_2 terminals each will be modelled by

$$T = \frac{P_1}{C-S} + \frac{P_2}{(C-S) - S/P_1} \quad (8)$$

The first term is the delay in the top level, consisting of P_1 dedicated subchannels. The second term is the delay in each one of the bottom-level subsystems, each of which is carrying a traffic of S/P_1 over P_2 dedicated subchannels. $C-S$ is the capacity available to every one of the bottom-level systems.

The T of (8) will be minimal when

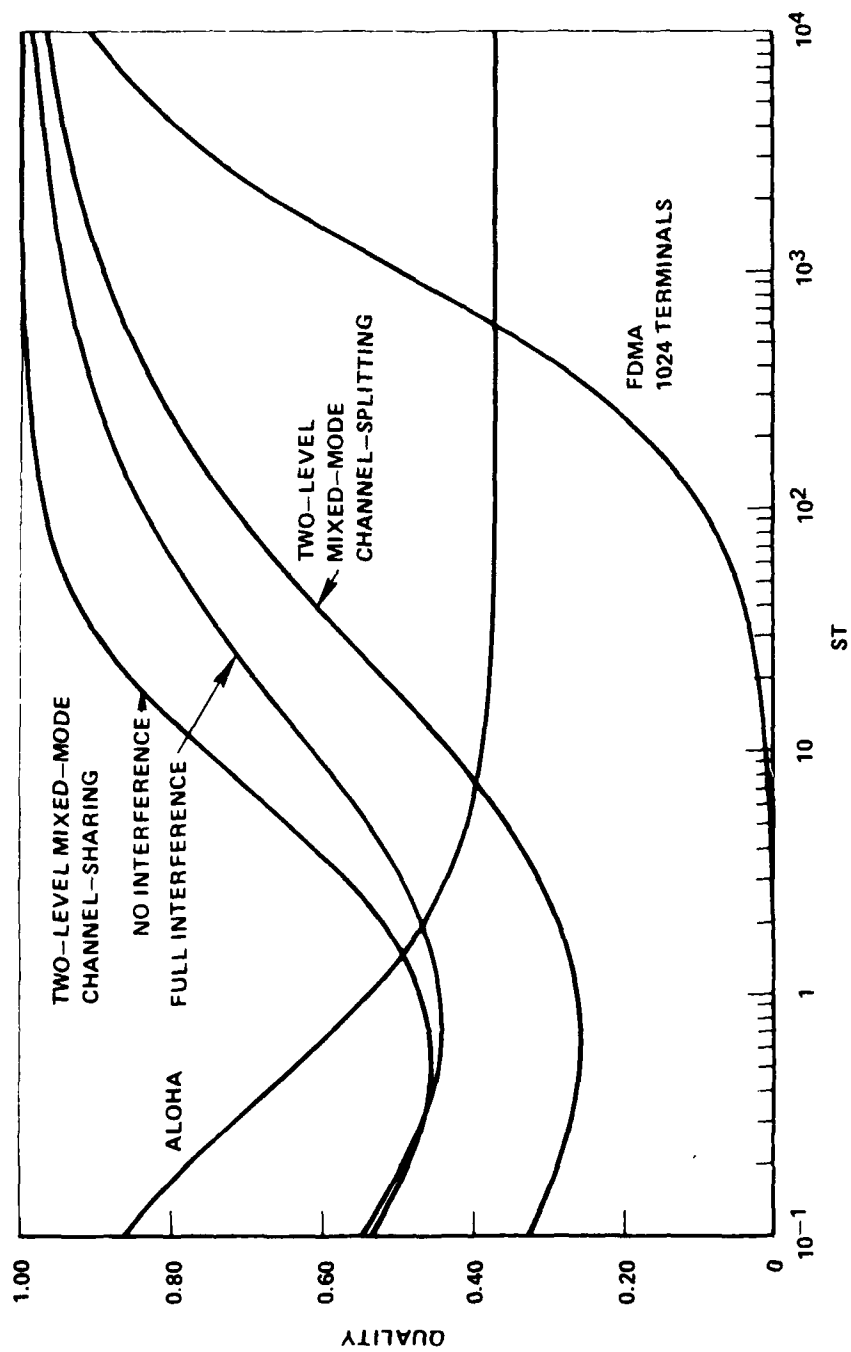


Figure 5. Sharing Versus Splitting in Two-Level Mixed-Mode Systems.

If a transmission from a terminal can be heard by more than one repeater, the system performance may be improved by allowing any of the repeaters, which received this transmission correctly, to relay it to the station [10]. In that case the message will not have to be correctly received specifically by its repeater, and success in reaching any of the repeaters will be enough. This advantage should be traded against the possibility that a message will successfully reach more than one repeater, and that all these repeaters will send it on. We shall not analyze this idea in any more detail.

5. Sharing or Splitting?

In the previous sections we have introduced several models for a two-level system in which both levels share the communication channel. But is this sharing good? In order to answer this question, consider another alternative:

If we have a communication medium with capacity C , let us assign a portion βC to the terminal-repeater traffic and a portion $(1-\beta)C$ to the repeater-station traffic. Using R repeaters, and assuming no interaction among ALOHA subsystems, we get the following equation for T :

$$T = \frac{R}{\beta C - S} + \frac{1}{(1-\beta)C - eS/R}$$

We can now minimize T by choosing both R and β . The minimum T will be obtained when the following two equations are satisfied:

$$R = \frac{eS}{\beta C - S}$$

$$\frac{\beta C - S}{C(1-\beta)} = \frac{\sqrt{eS}}{\sqrt{eS} + \sqrt{\beta C - S}}$$

These equations can be solved numerically, and Fig. 5 gives the quality of this optimal two-level channel-splitting organization compared with ALOHA, FDMA with 1024 terminals. Also included are the channel sharing scheme, in the cases of no interference and full interference. We see that sharing the channel is significantly better than splitting it.

Sharing is superior to splitting in very general circumstances, as the following theorem shows:

Theorem 1: Consider a two-level terminal-station communication system using a broadcast channel. This channel can either be split between levels or shared by both. Assume the channel-sharing mode has the following two properties:

- (1) Top-level communication is not bothered at all by bottom-level communication, i.e., the bottom level is transparent
- (2) The only effect activity in the top level has on the bottom level is to subtract itself from the capacity available to the bottom level

Then the channel sharing mode is superior to the channel splitting mode.

Proof: Let us start with a channel-splitting system carrying a given traffic and modify it to get a channel-sharing system that will carry the same traffic with a smaller delay. When the new top level is active it uses all the available bandwidth. Its transmission time will therefore be shorter than the transmission time in the channel-splitting system. By appropriate scaling and adjustment of the transmission policy in the top level we can ensure it will have an equal or shorter waiting time, and that it will utilize the same fraction of the total communication resource as did the old top level. The delay in the new top level will therefore be smaller than the delay in the old top level. Since the old top level must have been less than fully utilized, some of the capacity assigned to it in the channel-splitting

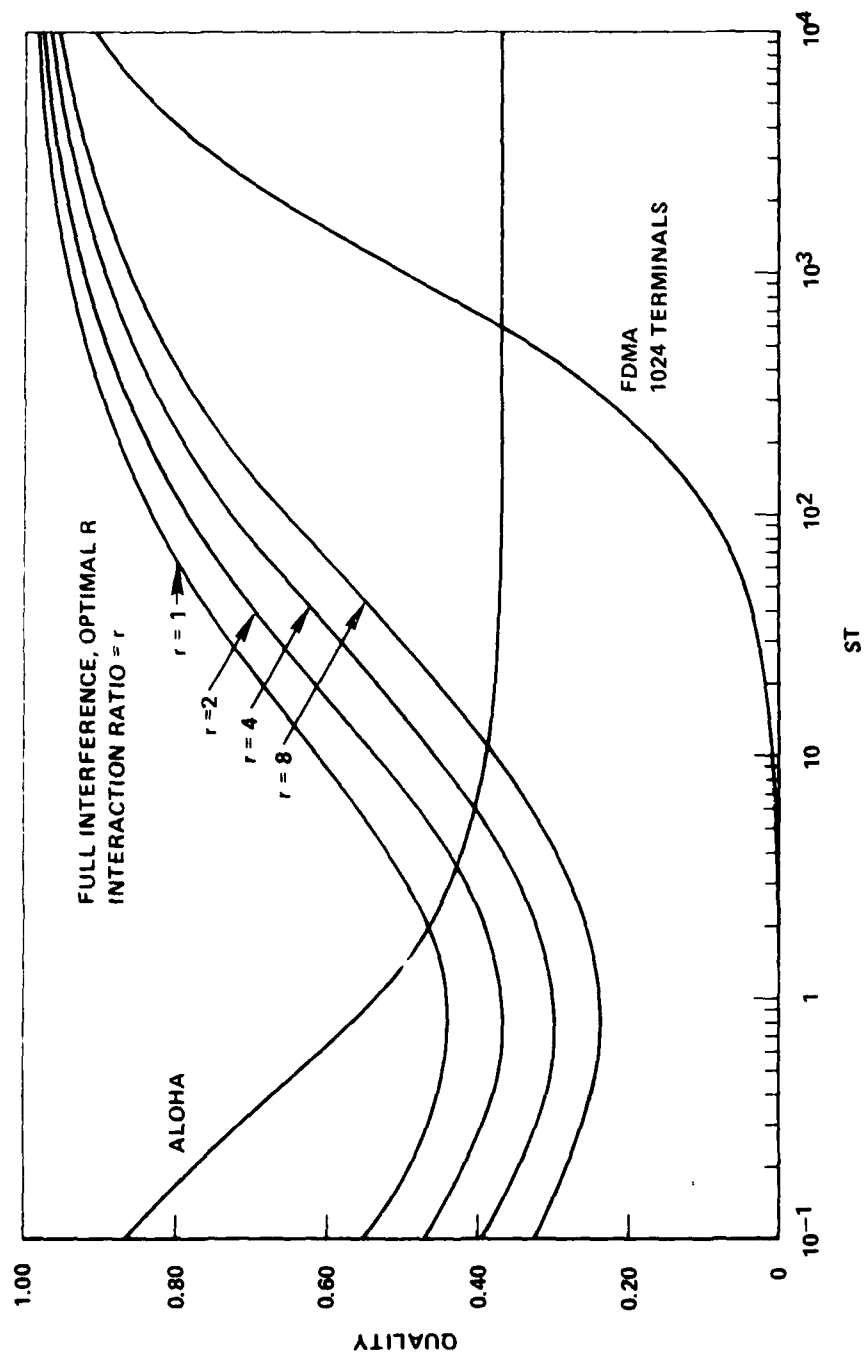


Figure 4. Two-Level Mixed-Mode With Interacting ALOHA Subsystems.

Let us stress once again that we do not explicitly treat the question of transmission errors. Instead of discussing the probability of successful reception and its dependence on various parameters we use the following simple model: A transmission is always received correctly if its source is within range of the destination and if there is no interference at the destination. Interference is caused by any other transmission within range of the destination. The range dependence of a successful reception is modelled as a step-function. When there is no interference, a transmission will always be successful if the distance to the destination is less than the range, and will never be successful if the distance is more than the range.

Let A_1 be the area covered by the any group of terminals, intended to be heard by one repeater. Let A_2 be the area covered by those terminals which are actually heard by the repeater. In any safe design we must have $A_2 > A_1$. Let r be the ratio A_2/A_1 . r will obviously depend on the shape of the cells around the repeater and on the terminal's power. What is the effect of the number of repeaters on r ? A simple geometric argument leads to the following conclusion: If we change the number of repeaters and the size of their cells, hold fixed the shape of the cells and adjust the terminal's power to get the same power at the repeater in the worst case (which is when the terminal is as far as it can be from the nearest repeater) then r will stay the same.

For example, take the case of a plane divided into identical regular hexagons. Let us give every terminal exactly the power necessary, on the average, to reach the center of a hexagon from its vertices, without any margin of safety. In this case r will be the ratio between the area of a circle and the area of an inscribed regular hexagon, i.e., $r=1.209$. If we wish to guarantee that each terminal can reach more than one repeater the transmission range must be equal to the (worst case) inter-repeater distance. In this case r will be equal to 3.627.

For a given shape of cell and power adjustment policy, we have therefore a set of interacting ALOHA systems, where the amount of interaction does not depend on the number of repeaters. A simple argument, like that used to find the maximum utilization of slotted ALOHA system [9], leads to the following: The maximum utilization of each ALOHA system consisting of a repeater and its terminals will be degraded by the interference of its neighbors, and is equal to $1/re$.

Modifying (4) we get for our present two-level system

$$T = \frac{R}{C-S} + \frac{1}{(C-S) - erS/R} \quad (6)$$

The optimal R is now given by

$$R_{\text{optimal}} = \frac{1}{e} \left[\frac{C-S}{erS} + \sqrt{\frac{C-S}{erS}} \right]$$

and T with this optimal R is

$$T = \frac{1}{C-S} \left[\sqrt{\frac{C-S}{erS}} + 1 \right]^2$$

Fig. 4 shows the quality of the 'full interference' case when interaction among different ALOHA systems exists. r , the coefficient of interaction, takes there the values 1, 2, 4 and 8.

In general, with more interaction, we shall be able to achieve a lesser portion of the infeasible region, and more repeaters will be needed. But having neglected the cost of repeaters, we should certainly not allow their number to grow without limit. Another problem with large R is that we have assumed that the terminal population is infinite. But when R becomes comparable to our actual number of terminals, the one-level FDMA will, of course, be better than this two-level organization.

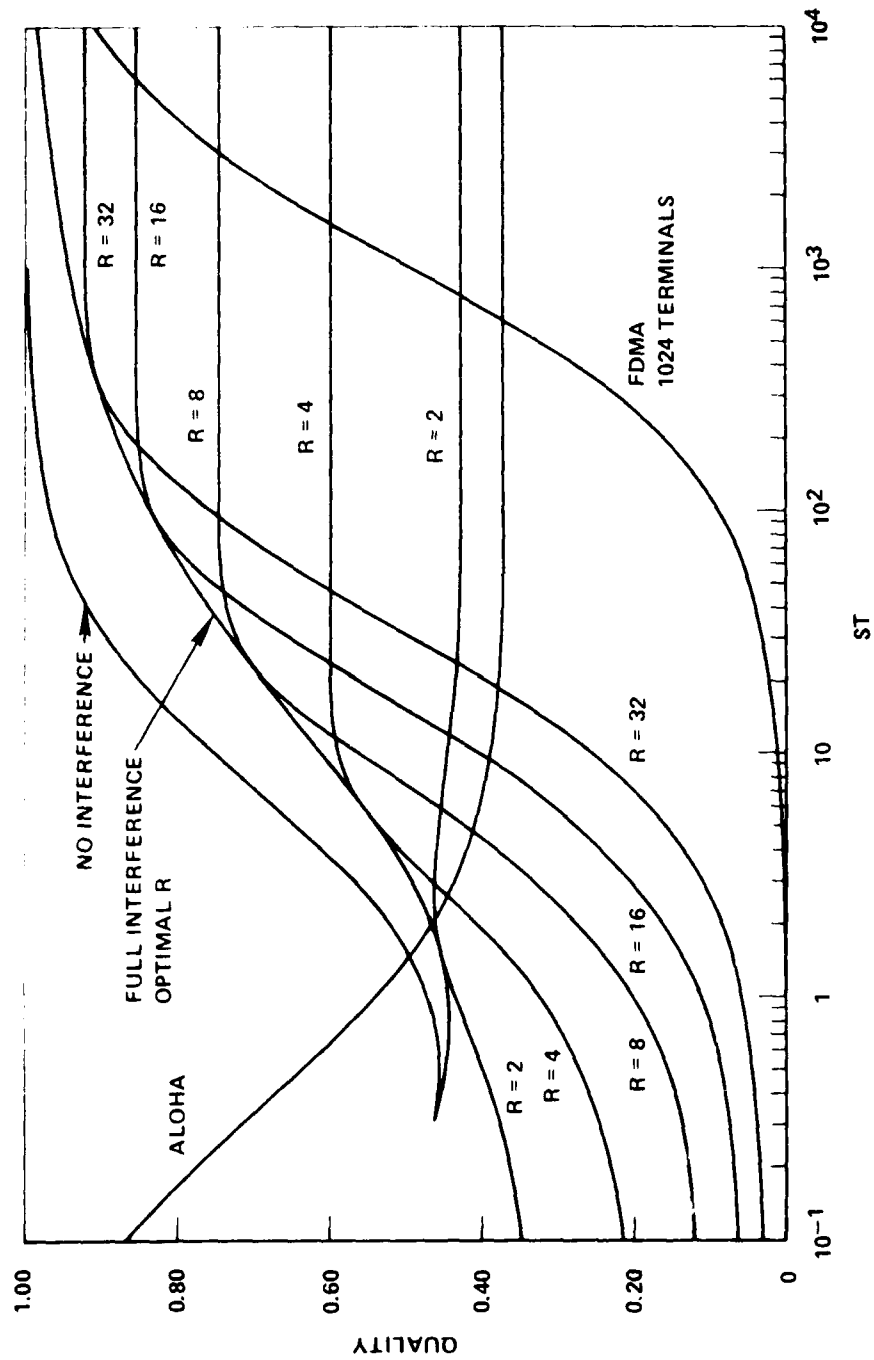


Figure 3. The Two-Level Mixed-Mode "Full Interference" System.

Even if the geometry alone is not enough to justify the assumption of transparent bottom, there are other good reasons to consider it valid. Since we expect to have few repeaters, they may be expensive and sophisticated. We shall assume now that repeaters are powerful and sophisticated enough to be perfectly captured by the station in the presence of bottom-level transmissions. The top level will never 'see' the bottom level, and this is the reason for the name 'transparent bottom'.

The assumption of perfect capture answers some of the problems raised at the beginning of the section. To model the effect of the other problems, we shall modify the 'no interference' assumption and assume that a repeater cannot listen to its terminal whenever *any* of the repeaters is transmitting to the station. Calling this new assumption 'full interference' [8], we shall use it as a worst case estimate for the interference between repeaters and terminals. With the full interference assumption, the effective capacity available to each terminal group is $C-S$, and instead of (1) we have for T the following expression:

$$T = \frac{R}{C-S} + \frac{1}{(C-S) - eS/R} \quad (4)$$

The optimal R is given by

$$R_{\text{optimal}} = \frac{eS + \sqrt{eS(C-S)}}{C-S} \quad (5)$$

and T with this optimal R is given by

$$T = \frac{1}{C-S} \left[\sqrt{eS/(C-S)} + 1 \right]^2$$

Fig. 3 shows the quality of the two-level hierarchy under the 'full interference' assumption. A significant part of the 'infeasible' region is still filled, but many more repeaters are necessary in order to achieve this. From (5) we see that as $S \rightarrow C$, $R \rightarrow \infty$. The quality in Fig. 3 is given for optimal R , and for R fixed at 2, 4, 6, 8, 16 and 32. For comparison Fig. 3 also includes ALOHA, FDMA with 1024 terminals, and the two-level 'no interference' case of the previous section with optimal R . In both curves with optimal R only the portion with $R > 1$ is drawn. They start at the same point because when $R=1$ the 'no interference' and the 'full interference' assumptions are identical.

4. Interacting ALOHA Subsystems

Spatial reuse is another strong assumption made in section 2: each repeater will be heard by 'its' receiver and by no other receiver. Is this a reasonable requirement? We do not mind installing a few sophisticated repeaters but the many terminals should be cheap and simple. These terminals may be mobile or unattended and they will not necessarily know where they are or where their repeater may be. Even if each terminal had a directional antenna or an adjustable output power, it might not have the information necessary to control them. Let us assume that all terminals have the same power and an omnidirectional antenna.

Consider a division of the plane into a set of equal polygons. In the 'middle' of each we place a repeater. Assume the terminals are uniformly placed over the plane. We wish to guarantee that a terminal will be heard by its nearest repeater. If the only factor that determines reception is power at the receiver, we must give each terminal enough power for the worst case (when its distance to the nearest repeater is maximal). We shall assume that whenever two terminals have enough power to be heard by the same repeater, the resulting interference will destroy both messages, that is, there is no capture of the terminals' transmissions. Because every terminal is given enough power for the worst case range, some terminals will be heard by more than one repeater. The assumption of no interaction between terminal groups must, therefore, be modified.

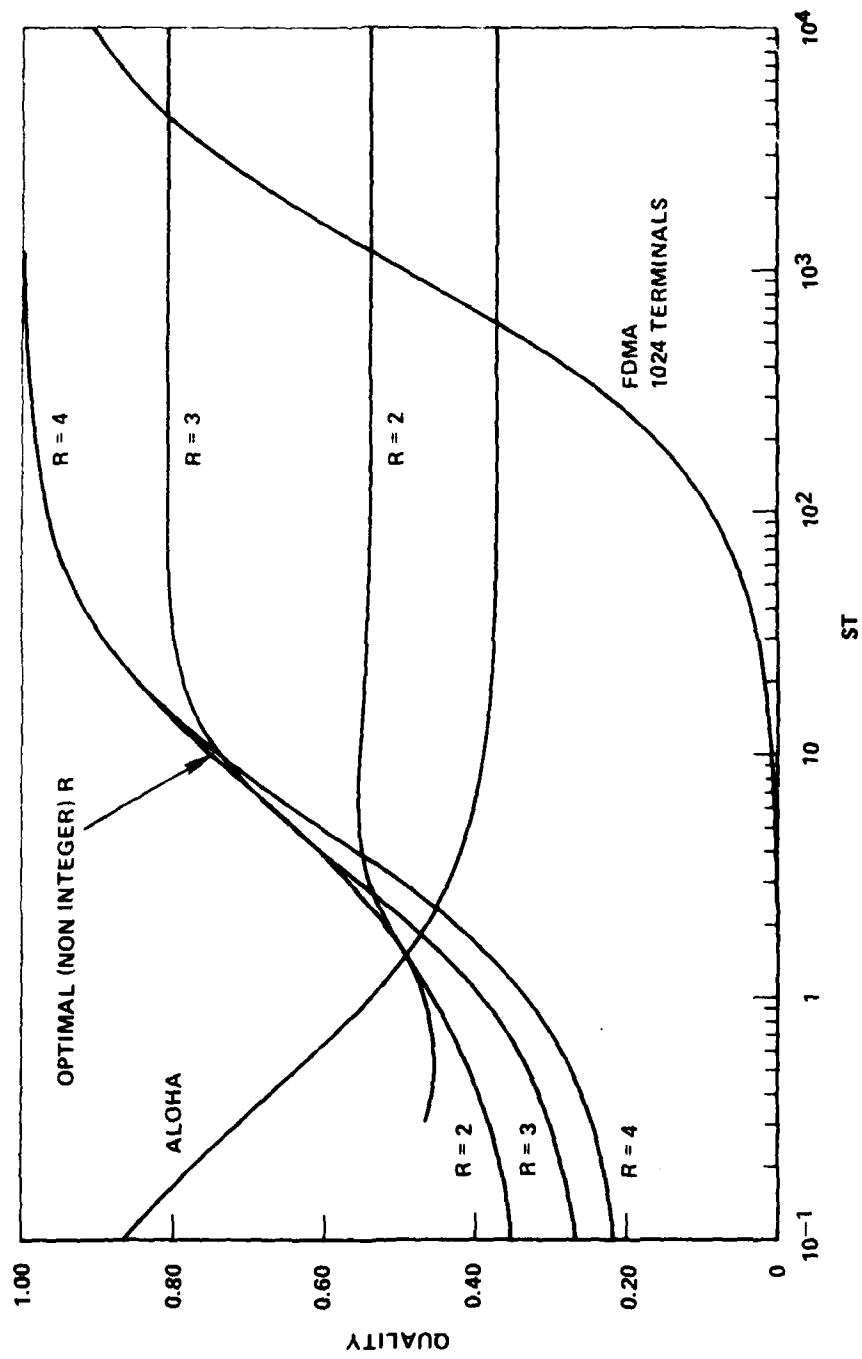


Figure 2. The Two-Level Mixed-Mode "No Interference" System.

use for delay in ALOHA systems.

When both levels are thus slotted and synchronized, the effective capacity available to each terminal group will be equal to $C - S/R$, since S/R of the available capacity is used by its upper level repeater. The load on each lower level ALOHA system accessing a given repeater will be S/R . The average time T a message will spend in the system is therefore

$$T = \frac{R}{C - S} + \frac{1}{(C - S/R) - eS/R} \quad (1)$$

where the first term is the time spent in the top level (repeater-station) and the second is the time spent in the bottom level (terminal-repeater).

With (1) giving the total time in system, we can now ask what is the optimal number of repeaters. Minimizing T we get

$$R_{\text{optimal}} = \frac{1}{C} \left\{ (1+e)S + \sqrt{(1+e)S(C-S)} \right\} \quad (2)$$

With this optimal R we get for T

$$T_{\text{minimal}} = \frac{1}{C} \left\{ 1 + \sqrt{(1+e)S/(C-S)} \right\}^2 \quad (3)$$

From (2) we can see that when S is very small, the optimal R is almost zero. This occurs because the two-level structure is worse than the one-level ALOHA when $S \rightarrow 0$. The optimized R will try to compensate for this by driving to zero the time spent in the top level. We can also get from (2) that the largest optimal R is 3.95, obtained when $S/C = .944$. In practice, R must be an integer greater than one.

Equation (3) gives T as a function of S and C . The quality can be calculated by comparing C with the capacity necessary in an M/M/1 scheme for the same S and T . That is, $Q = (S+1/T)/C$. Fig. 2 gives the quality of the two-level structure with the optimal R (which is not necessarily an integer). The section of the curve in which the optimal R is smaller than 1 is not drawn. Also plotted is the quality of the two-level structure, when T is given by (1), and when R is fixed at 2, 3, and 4. For comparison, the figure also gives the quality of ALOHA and the quality of FDMA with $m=1024$ terminals.

We see that a two-level system can fill in a large portion of the 'chasm' left between ALOHA and FDMA. This chasm is an 'infeasible' region when only ALOHA and FDMA are considered. When the number of terminals grows, FDMA will move even further to the right, but ALOHA and our two-level scheme will not be modified (both of these already assume an infinite population of terminals), so the relative gain achieved by the two-level hierarchy over both ALOHA and FDMA will be even greater.

This seems almost too good to be true! In the following sections we shall reexamine our assumptions and see how relaxing them will modify and degrade the result.

3. The 'Full Interference' Case

Some strong assumptions were made in the last section to the effect that both terminals and repeaters can use the same broadcast channel, with minimal interference. Consider first the assumptions of 'transparent bottom' and 'no interference'. These assumptions are reasonable if all the terminals are far from the station, for example if they are spread around a ring with the station in the middle. But if there are terminals close to the station, more interference may occur. Transmissions from a terminal situated near the station to its repeater may interfere with repeater-station communication, and transmissions from one repeater to the station may interfere with transmissions from terminals to another repeater.

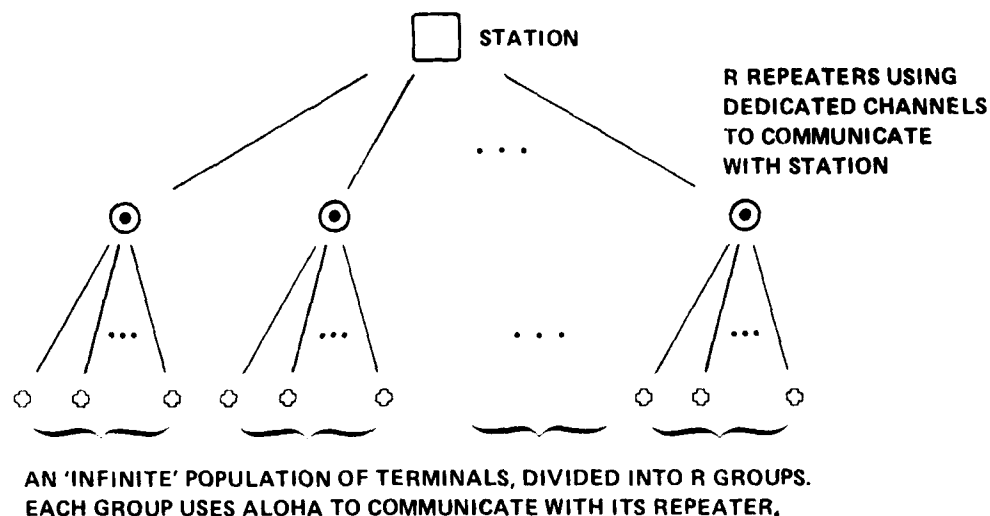


Figure 1. Two-Level Mixed-Mode Broadcast Systems.

- (2) *spatial reuse*: The terminal-repeater communication will be done using the ALOHA scheme. Each of the R groups can use the entire bandwidth to communicate with 'its' repeater and there will be no interference between transmissions of the terminals in different groups. That is, the terminals in each group will be heard by exactly one and the same repeater.
- (3) *transparent bottom*: Bottom-level transmissions have no influence on top-level transmissions. Each repeater will use a dedicated subchannel whose capacity is equal to $1/R$ of the *total* available capacity for its communication with the station.
- (4) A repeater cannot listen to its terminals while it is transmitting to the station.
- (5) *no interference*: A repeater's capacity to listen to its terminals will not be bothered by any of the other repeaters transmitting to the station [8].

The throughput-delay performance of the ALOHA schemes is not described by a simple analytic expression [2]. For simplicity we model the delay T in an 'infinite' population ALOHA system carrying a traffic S on a channel whose bandwidth is C by $T = \frac{1}{C - eS}$. This is a simple two-parameter approximation, that reproduces the known behavior of (unslotted) ALOHA when $S=0$, and the known behavior of (slotted) ALOHA when $S/C=1/e$. For a similar three-parameter approximation see [5].

In our two-level scheme, if a terminal is trying to transmit to its repeater while the repeater is transmitting to the station, the terminal will not be successful, and will have to try again. To minimize the wasteful effect of these bottom-level failures the two levels should be slotted and synchronized. This means that dedicating subchannels in the top level must be done by Time Division Multiple Access (TDMA). Despite the fact that TDMA must be used, we shall describe the delay in the top level by the FDMA formula, which is both simpler and more similar to the M/M/1 type formula we

favorable conditions. In sections 3 and 4 we relax the assumptions on the interaction between levels and on the interaction between lower-level ALOHA groups respectively. Section 5 shows that sharing the channel by both levels is often better than splitting it. Mixing dedicated channels and any random access scheme is discussed in section 6, and section 7 shows that having more than two levels leads to only a small improvement.

In the second half we discuss systems in which distributed terminals are both the sources and destinations of all messages. We shall call such systems *networks*, and assume that the average distance travelled by messages is much larger than the distance between a terminal and its nearest neighbors. Section 8 introduces two-level mixed-mode networks with the simplest possible routing. Section 9 shows that improving the random access level leads to a relatively small overall improvement, and section 10 similarly shows that introducing more than one dedicated level leads to a small improvement.

Throughout the paper we assume that the communication resource available is a *broadcast* channel of capacity C . We shall also assume that the message arrival process is Poisson with a total rate S ; that message lengths have an exponential distribution; and that all terminals contribute equally to the overall traffic. This last assumption characterizes the case which is hardest to control efficiently. We choose the data unit so that the average length of a message is equal to 1. This is simply a convenient normalization, which is equivalent to measuring communication capacity in messages (of an average length) per second, instead of measuring in bits per second.

If the terminals were co-located in the same place, the best access scheme would be to form a queue of busy terminals (i.e., those having anything to transmit) and to let them use the full bandwidth available one after the other. Forming one queue is much better than giving each terminal a fraction of the bandwidth, and letting each terminal queue its own messages [6]. When terminals are distributed and cannot form one queue without some investment in coordination and control more bandwidth will be necessary. Assuming that S and T are given, we define the *quality* [7] of an arbitrary access scheme as the inverse ratio between the capacity necessary when using this scheme and the capacity necessary when using the *best possible* scheme, in which messages form one queue and share one channel. When messages arrive independently and their lengths are exponentially distributed, this best scheme is the $M/M/1$ queue, in which we have $C_{M/M/1} = S + 1/T$.

2. The 'No Interference' Case

Given our broadcast channel, let us build a two-level *hierarchical* system by dividing the large number of terminals into R equal groups, and by giving each group a repeater. Each message will go from its terminal to its repeater, and from the repeater to the station. The terminal-repeater (bottom) level will have a large terminal population, possibly bursty, while the number of repeaters will, hopefully, be small, with enough traffic going through each for the repeater-station (top) level to be steady. It is natural, therefore, to suggest using ALOHA for the terminal-repeater level, and using dedicated channels for the repeater-station level.

Using ALOHA for the bottom level is desirable for other reasons too. For example, because no explicit control is exercised over transmission, ALOHA is especially good for mobile terminals and for situations where the number of potentially active terminals is much greater than the actual number active at any moment.

In order to model this two-level mixed mode centralized system, shown in Fig. 1, we shall start with the following assumptions (the words in *italics* will serve as names for the assumptions):

- (1) *channel sharing*: The communication medium is a broadcast channel, and both levels (terminal-repeater and repeater-station) may use the full bandwidth.

On the Advantage of Mixing ALOHA and Dedicated Channels

Abstract

When many terminals which are distributed in space must share communication resources, we face the following problem: What scheme can control the access to the communication resources in an effective way? We shall assume that S , the traffic to be carried, and T , the acceptable average delay, are specified, and that the goal is to design the least cost system satisfying these specifications.

Dedicating a fraction of the resources to some source-destination pairs is one very simple access scheme. Another simple scheme is ALOHA. When we combine the specified traffic and delay into the dimensionless quantity ST , whose inverse we call *burstiness*, we find the following: Dedicating separate channels is good when the traffic is steady, but bad when the traffic is bursty. ALOHA is good when the traffic is bursty, but bad when the traffic is steady. Neither ALOHA nor dedicated channels are good when the traffic is of medium burstiness.

Mixed-mode systems, using ALOHA in a bottom level and dedicated channels in a top level, can be good, since they can trade the amount of interference in the random access level against the number of dedicated channels in the top level. By choosing the right mix, such networks can become insensitive to the limitations of both access schemes.

1. Introduction

When many terminals which are distributed in space must share communication resources, we face the following problem: What scheme can control the access to the communication resources in an effective way? We shall assume that S , the traffic to be carried, and T , the acceptable average delay, are specified, and that the goal is to design the least cost system satisfying these specifications. Furthermore, we shall assume that only the capacity, i.e., bandwidth, necessary has a cost, and that equipment and transmission power are free.

Dedicating a portion of the resource to source-destination pairs is one very simple access scheme. Another simple scheme is ALOHA [1,2]. When we combine the specified traffic and delay into the dimensionless quantity ST , we find the following: The dedicated-channel scheme is good when $ST \gg 1$ (the traffic is then said to be *steady*) but bad when $ST \ll 1$ (the traffic is then said to be *bursty*). ALOHA is good when the traffic is bursty, but bad when the traffic is steady. Neither ALOHA nor dedicated channels are good when the traffic is of medium burstiness.

It is possible to improve the dedicated channel scheme when the traffic is bursty by a hierarchical structure that makes sharing of few high capacity channels possible [3]. It is also possible to improve the ALOHA scheme when the traffic is steady by trading off transmission range and the necessary number of hops [4]. Is it possible to obtain a good access scheme for medium burstiness by mixing the dedicated-channels and the ALOHA schemes? Kleinrock [5] has shown that splitting the resources and the traffic between two access schemes can never lead to an improvement. Here we show that by building a *hierarchical* system with different schemes used at different levels we can get a significant improvement at medium burstiness. The first half of this paper applies this idea to systems in which the sources of messages are many terminals distributed in space, but in which all messages are destined to one common *station*. We shall call such systems *centralized*, and assume that m , the number of terminals, is very large. In section 2 we introduce the mixed-mode scheme under the most

- [9] Kleinrock, L., "Performance of Distributed Multi-Access Computer Communication Systems," *Information Processing 1977, Proceedings of IFIP Congress 77*, Toronto, August 1977, North Holland, Amsterdam, 1977, pp. 547-552.
- [10] Abramson, N., "Packet Switching with Satellites," *AFIPS Conference Proceedings*, 1973 National Computer Conference, Vol. 42, pp. 695-702.
- [11] Kleinrock, L., and F.A. Tobagi, "Packet Switching in Radio Channels: Part I--Carrier Sense Multiple-Access Modes and their Throughput-Delay Characteristics," *IEEE Transactions on Communications*, Vol. COM-23, pp. 1400-1416, December 1975.
- [12] Meister, B., H. Muller, and H. Rudin, "New Optimization Criteria for Message-Switching Networks," *IEEE Transaction on Communication Technology*, Vol. COM-19, pp. 256-260, June 1971.
- [13] Metcalfe, R.M., and David R. Boggs, "Ethernet: Distributed Packet Switching for Local Computer Networks," *Communication of the ACM*, Vol. 19 pp. 395-404, July 1976.
- [14] Kahn, R.E., "The Organization of Computer Resources into a Packet Radio Network," *IEEE Transactions on Communications*, Vol. COM-25, pp. 169-178, January 1977.
- [15] Kamoun, F., "Design Considerations for Large Computer Communication Networks," Computer Systems Modeling and Analysis Group, School of Engineering and Applied Science, University of California, Los Angeles. UCLA-ENG-7642, April 1976. (Also published as a Ph.D. Dissertation, Computer Science Department)
- [16] Kleinrock, L., *Communication Nets: Stochastic Message Flow and Delay*, McGraw-Hill, New York, 1964. Out of print. Reprinted by Dover Publications, New York, 1972.
- [17] Coxeter, H.S.M., *Introduction to Geometry*, second edition, John Wiley and Sons, New York, 1969.
- [18] Newell, G.F., "Optimal Network Geometry," *Transportation and Traffic Theory, Proceedings of the Sixth International Symposium on Transportation and Traffic Theory*, Sydney, August 1974, Elsevier, New York, 1974, pp. 561-580.
- [19] Akavia, G.Y., and L. Kleinrock, "On the Advantage of Mixing ALOHA and Dedicated Channels," submitted for publication.

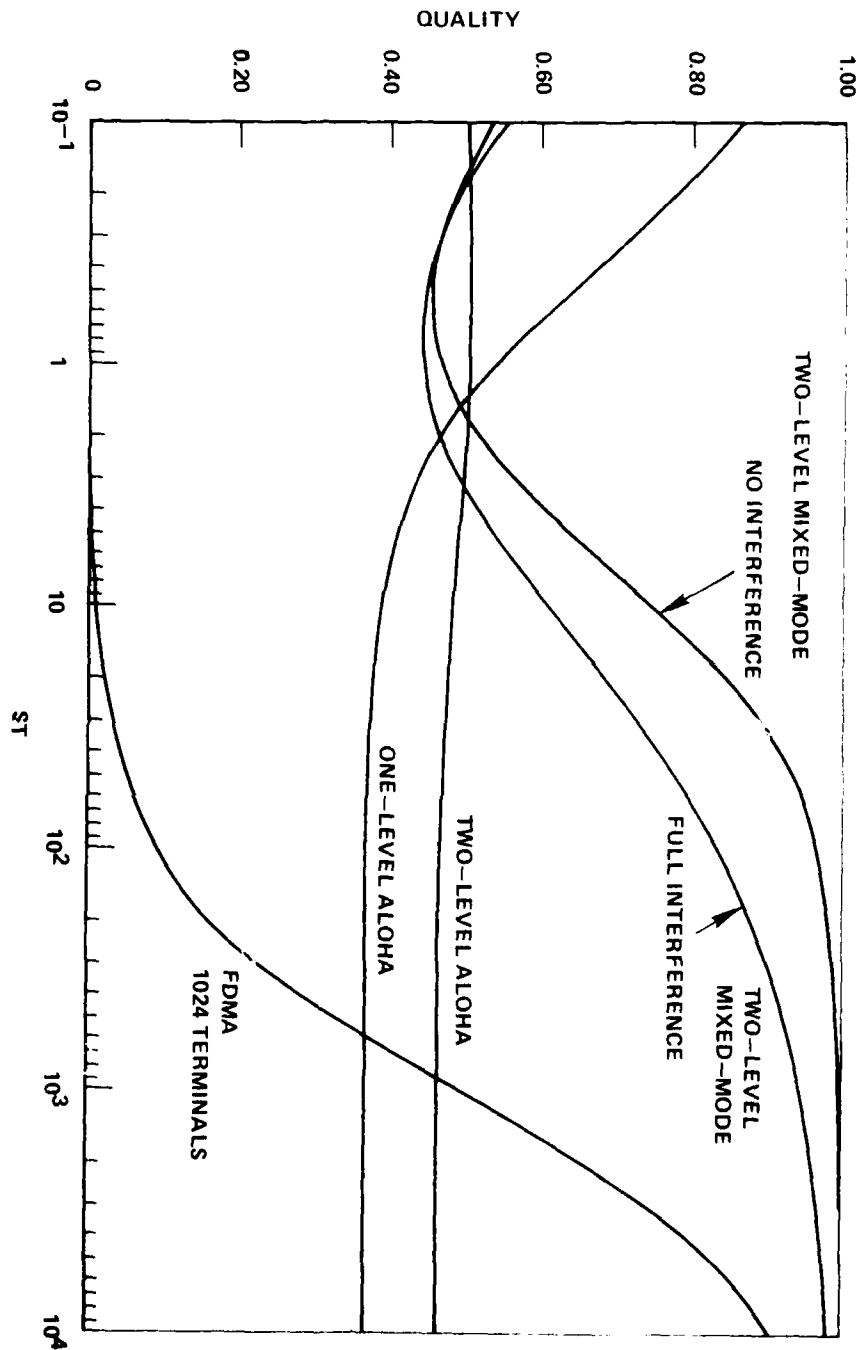


Figure 6. Mixed-Mode and Pure ALOHA Two-Level Organizations.

there still will be a region that is infeasible if we consider only FDMA and the given random access scheme, and the two-level mixed-mode scheme can help fill this infeasible region.

Let us divide the gap between .367 (the maximum utilization of ALOHA) and 1 (the maximum utilization of M/M/1) into four equal parts, and consider general random access schemes where U (the maximum utilization) is equal to .526, .684 and .842. Figures 7, 8 and 9 show the quality of the one-level and the two-level schemes, with this set of values for U . The mixed-mode curves were obtained from the formulas of this chapter by substituting $1/U$ for e . The two-level random access curve was obtained as follows:

Let us assume that the total offered traffic G and the throughput S are related, in a general random access scheme with m terminals, by

$$\frac{S}{C} = \frac{G}{C} \left(1 - \frac{aG}{mC} \right)^{m-1} \quad (10)$$

The maximum utilization (i.e., the maximum S/C) of this system will be obtained when $aG/C=1$, and is equal to

$$\frac{1}{a} \left(1 - \frac{1}{m} \right)^{m-1} \quad (11)$$

Equation (11) has its maximal value when $m=2$, and the best two-level system will therefore, once again, have two repeaters. Since we have denoted the maximum utilization of an 'infinite population' system by U we must have $a = 1/Ue$. In analogy to (4.4) we can, therefore, model the delay in a two-level mixed-mode system by

$$T = \frac{1}{C - 2aS} + \frac{1}{(C - G/2) - S/2U}$$

The first term is the delay in the repeater-station level, which has a maximum utilization of $1/2a$, as obtained from (11). The second term is the delay in each one of the terminal-repeater subsystems, where G is given in terms of S and C by (10) with $m=2$. (This very simple model for a two-level random-access system should not be applied when $U > 2/e = .736$, because the calculated maximum utilization of a two-terminal system will then be greater than one!)

From Figures 7, 8 and 9 we see that the conclusion formulated earlier for ALOHA systems actually applies to random access systems in general: two-level mixed-mode systems fill a significant part of the infeasible region. While our model for a system with two levels of random access may be considered too crude, it seems to say that two levels of *random access* do not offer a significant improvement, and are almost dominated by the two-level *mixed-mode* systems.

7. Are Three Levels Ever Necessary?

If two-level mixed-mode systems are good, would three-level systems be better? Consider, for example, a system consisting of one ALOHA level as the bottom level, and two dedicated levels on top.

Despite the fact that every message takes two hops in the dedicated levels we shall assume that only one hop, the longer one, influences other repeaters, and for this influence adopt the 'full interference' assumption. If the two dedicated levels do not share bandwidth, but the bottom level shares with both of them, we can write for the delay in this three-level system

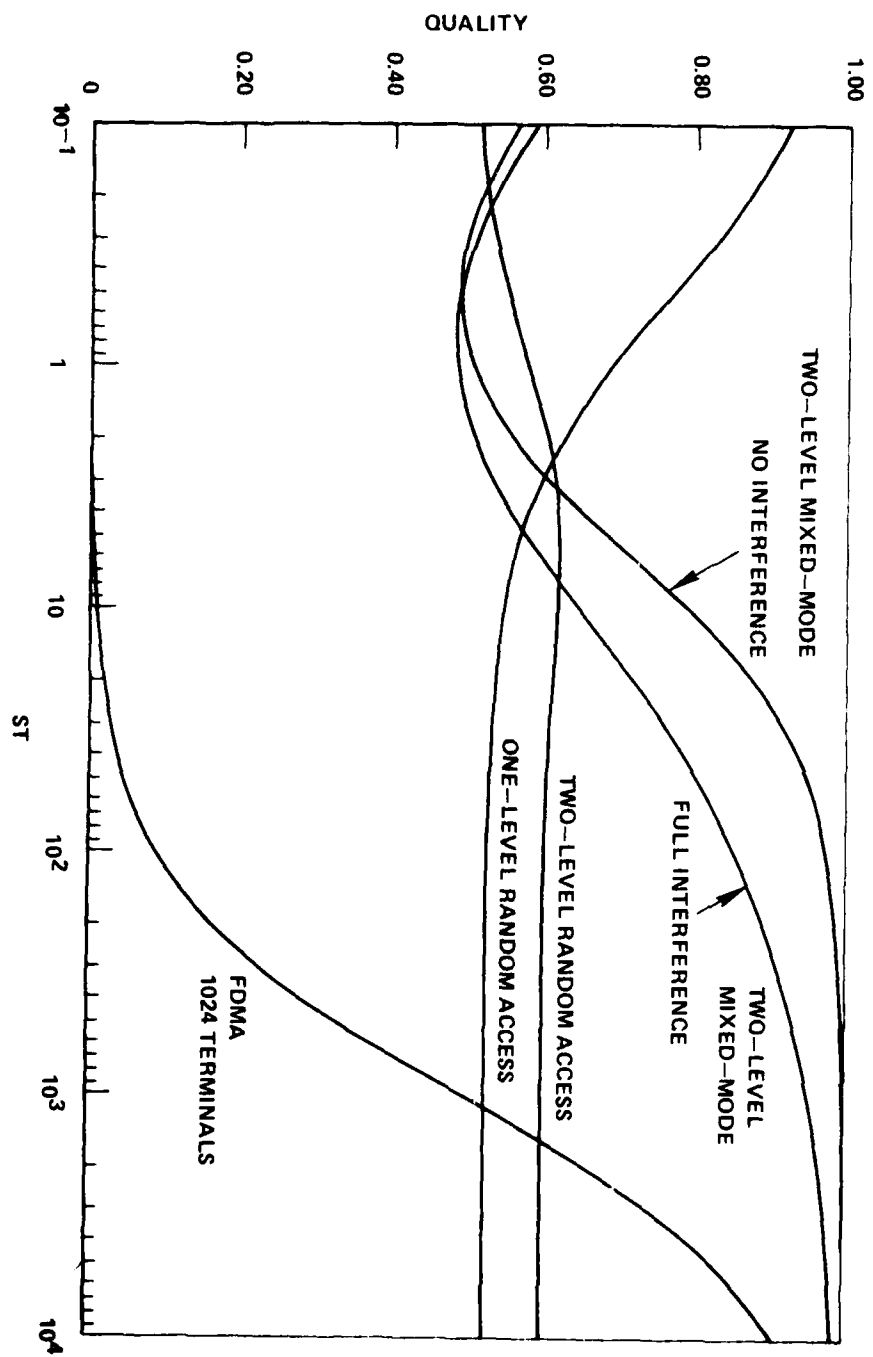


Figure 7. General Random Access, Maximum Utilization = 0.526.

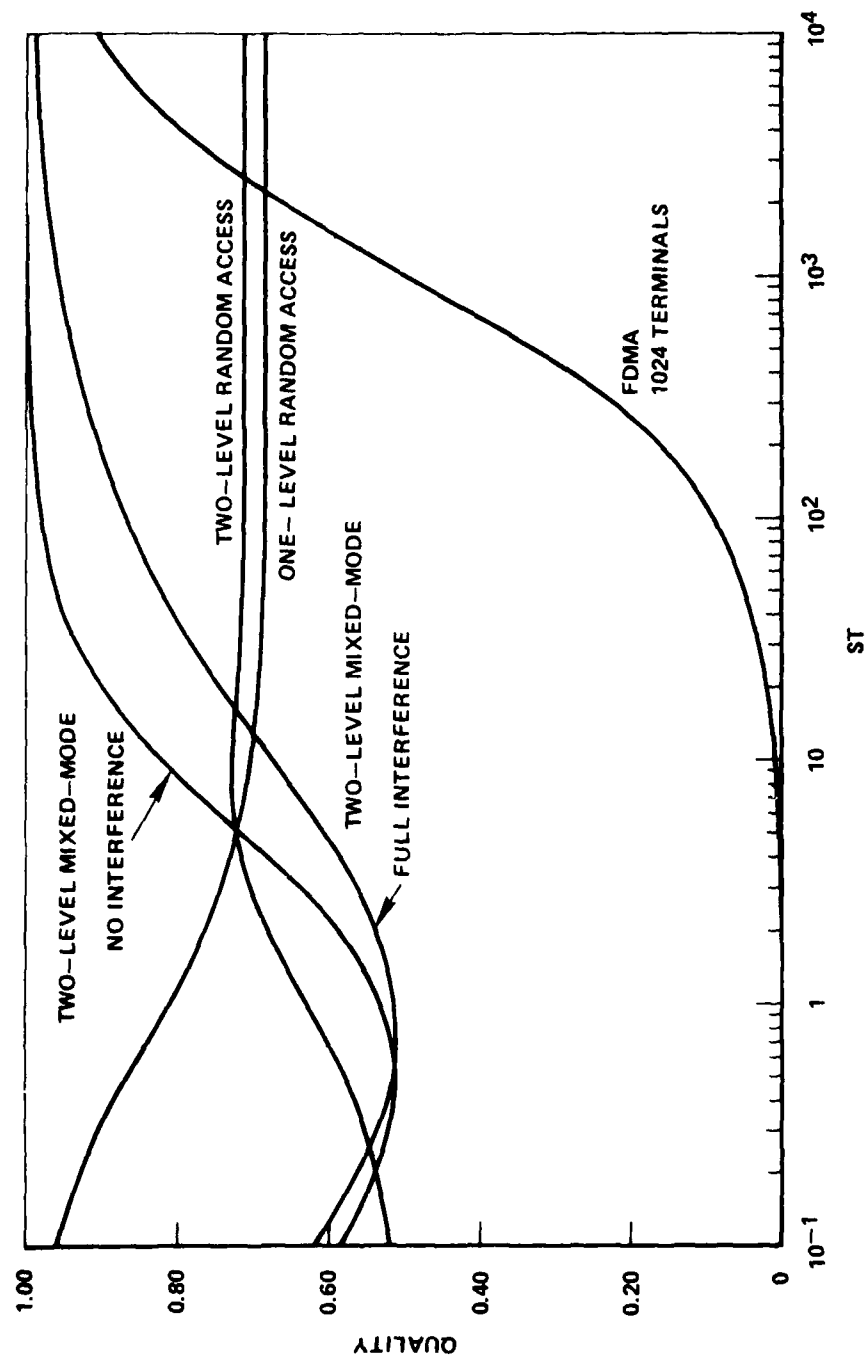


Figure 8. General Random Access, Maximum Utilization = 0.694.

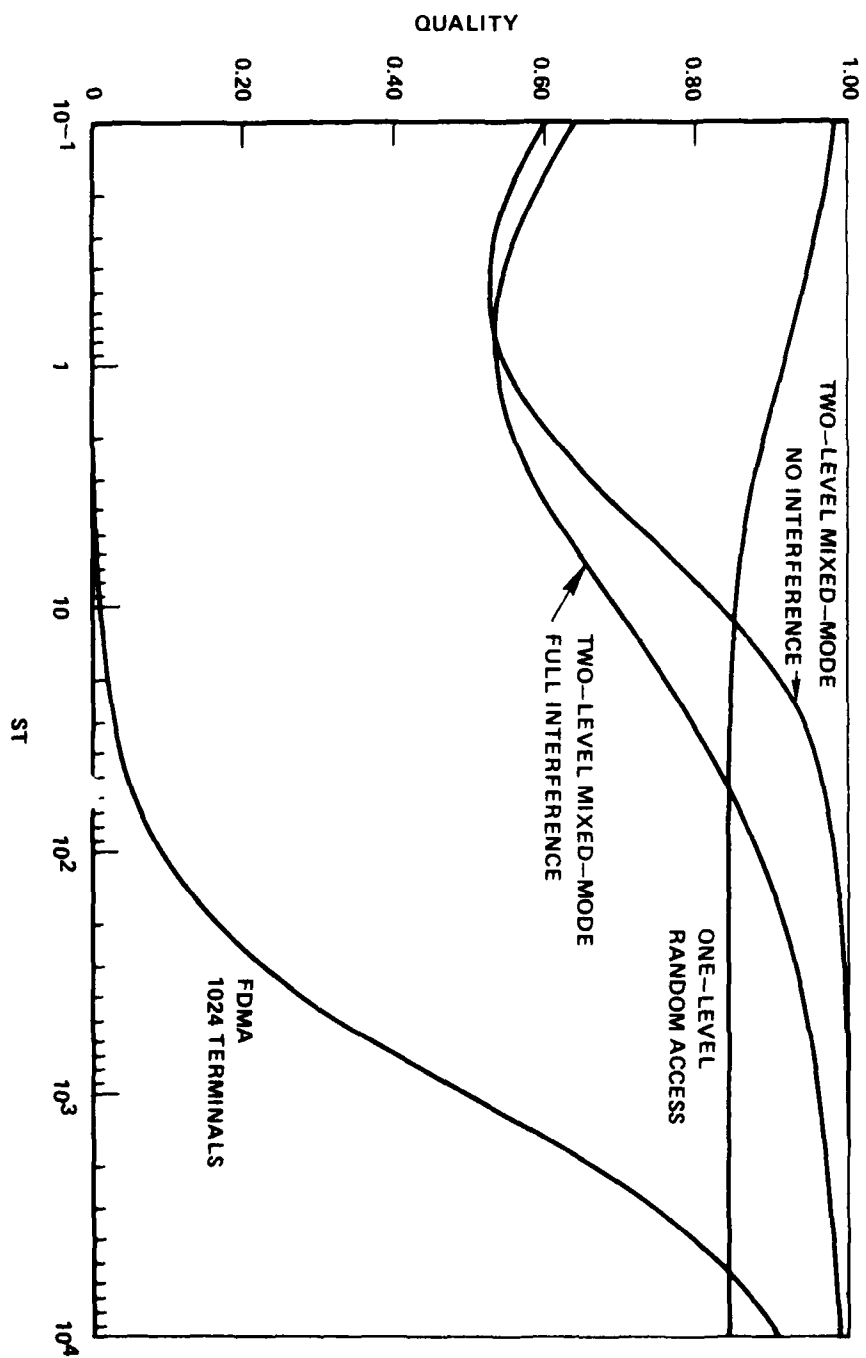


Figure 9. General Random Access, Maximum Utilization = 0.842.

$$T = \frac{4S}{(C-S)^2} + \frac{4\sqrt{R}}{(C-S)} + \frac{1}{C-S - erS/R} \quad (12)$$

The first two terms are the delay in the dedicated levels when we have R repeaters, obtained from (7). The third term is the delay in the ALOHA level, and r is the interaction ratio.

If we assume that the dedicated levels share the channel we can use for them (9), and the delay in this three-level organization is

$$T = \frac{S}{(C-S)^2} + \frac{2\sqrt{R}}{C-S} + \frac{1}{(C-S) - erS/R} \quad (13)$$

For a given C and S we can, in both (12) and (13), search for R , the number of repeaters, that will minimize T .

Fig. 10 shows the quality of the two-level and the three-level mixed-mode schemes, when there is no interaction between ALOHA subgroups (i.e., $r=1$) and the optimal number of repeaters was chosen in each as a function of burstiness. The three-level scheme was drawn only when it is better than the two-level scheme. Having three levels results in no noticeable improvement if the two dedicated levels split the channel and results in a small improvement if the two dedicated levels share the channel. The reason for this small improvement is clear: going from one dedicated level to two dedicated levels leads to a significant improvement only when the traffic is bursty and the number of repeaters is large. But in our two-level mixed-mode scheme the number of repeaters is large only when the traffic is steady, so adding a second dedicated level cannot lead to any dramatic improvement.

When we have interaction between the ALOHA groups, the number of repeaters becomes large earlier, i.e., when the traffic is bursty enough to make two dedicated levels better than one. Figures 11, 12 and 13 show the quality of the three-level mixed-mode scheme when the interaction ratio r is equal to 2, 4 and 8. The three-level scheme in which the channel is split between the two dedicated levels was actually plotted in all three figures, but becomes noticeable only when $r \geq 4$.

We see that introducing three levels improves the two-level performance significantly only when the interaction between ALOHA groups is very large. Even then, the gain achieved in going from two to three levels is much less than the gain achieved in going from one to two levels. When the interaction between ALOHA groups is strong, it may be unreasonable to ignore the interaction between repeater groups in the middle level. However, such an interaction-free division into groups was assumed in deriving (7) and (9), which form the basis for (12) and (13). Hence our three-level results are likely to be too optimistic. In reality, a three-level mixed-mode scheme will achieve an even smaller improvement over the corresponding two-level mixed-mode scheme than our figures show.

8. Two-Level Mixed-Mode Networks

In networks, i.e., when both sources and destinations are distributed, we have a situation similar to the one we saw earlier for centralized systems: it is easy to organize and to control (if any control is necessary) communication systems that are either very steady or very bursty, even if they are distributed. It is the distributed systems of medium burstiness that pose a problem. We saw earlier that a hierarchical two-level centralized system which mixes dedicated channels and ALOHA in the appropriate 'amounts' can be much better than either of them, for medium burstiness. Therefore, let us now apply the mixed-mode idea to networks. We shall discuss in detail only one-dimensional networks, but expect our major conclusions to be valid for two-dimensional networks too. Denote by N the average distance travelled by messages, and by S_d the rate of traffic originating in a unit length of the network.

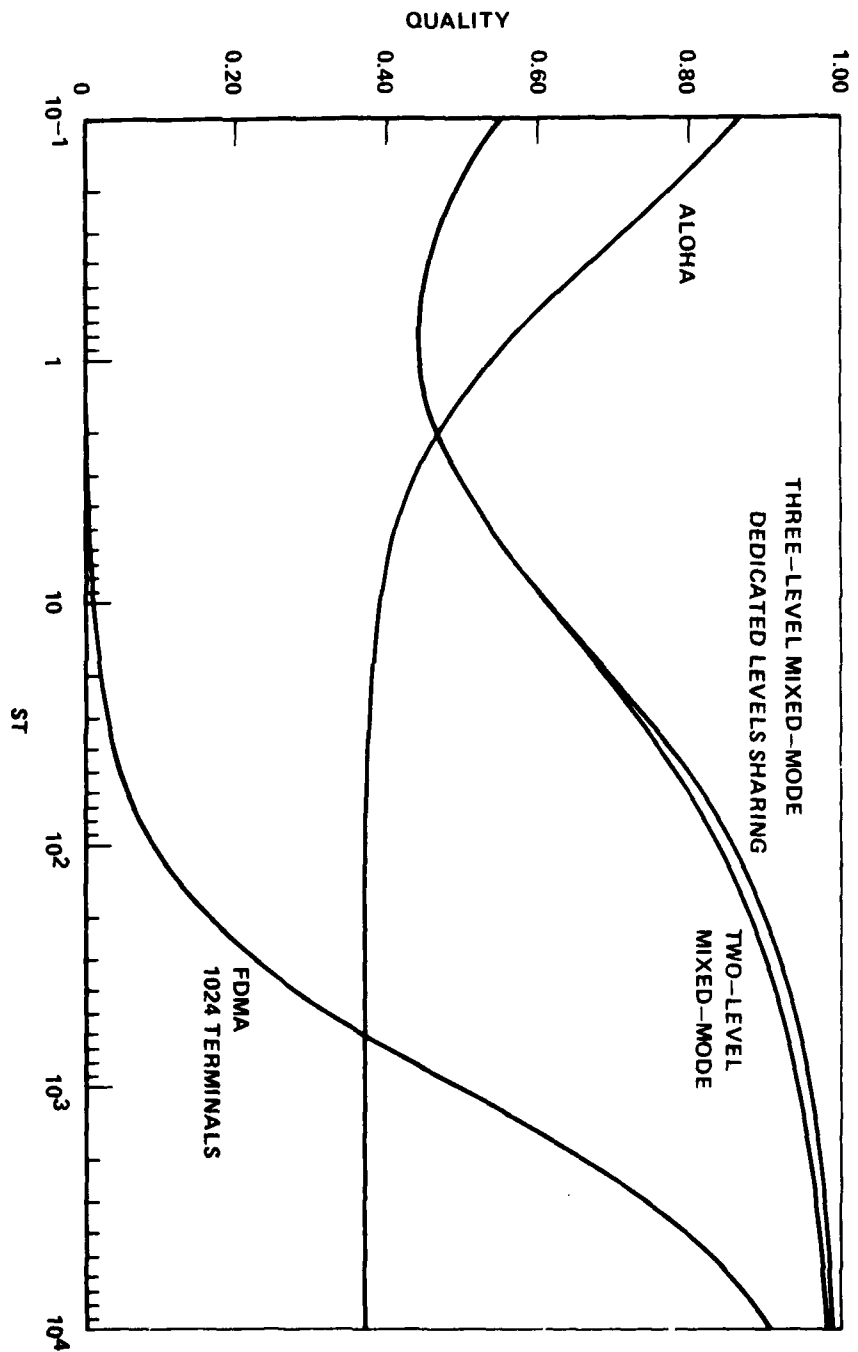


Figure 10. Quality of Three-Level Mixed-Mode Schemes,
No Interaction Between ALOHA Subgroups.

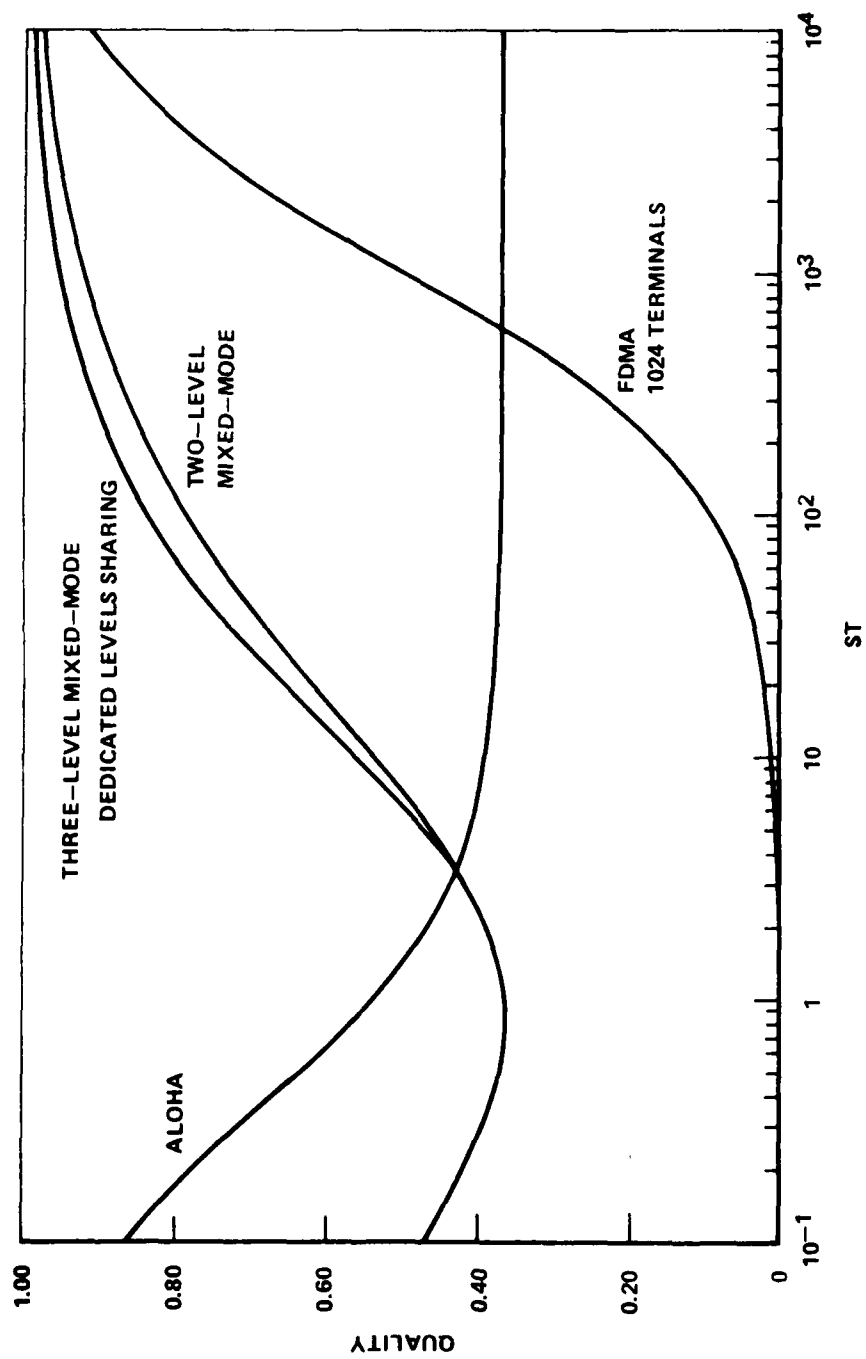


Figure 11. Quality of Three-Level Mixed-Mode Schemes, Interaction Ratio = 2.

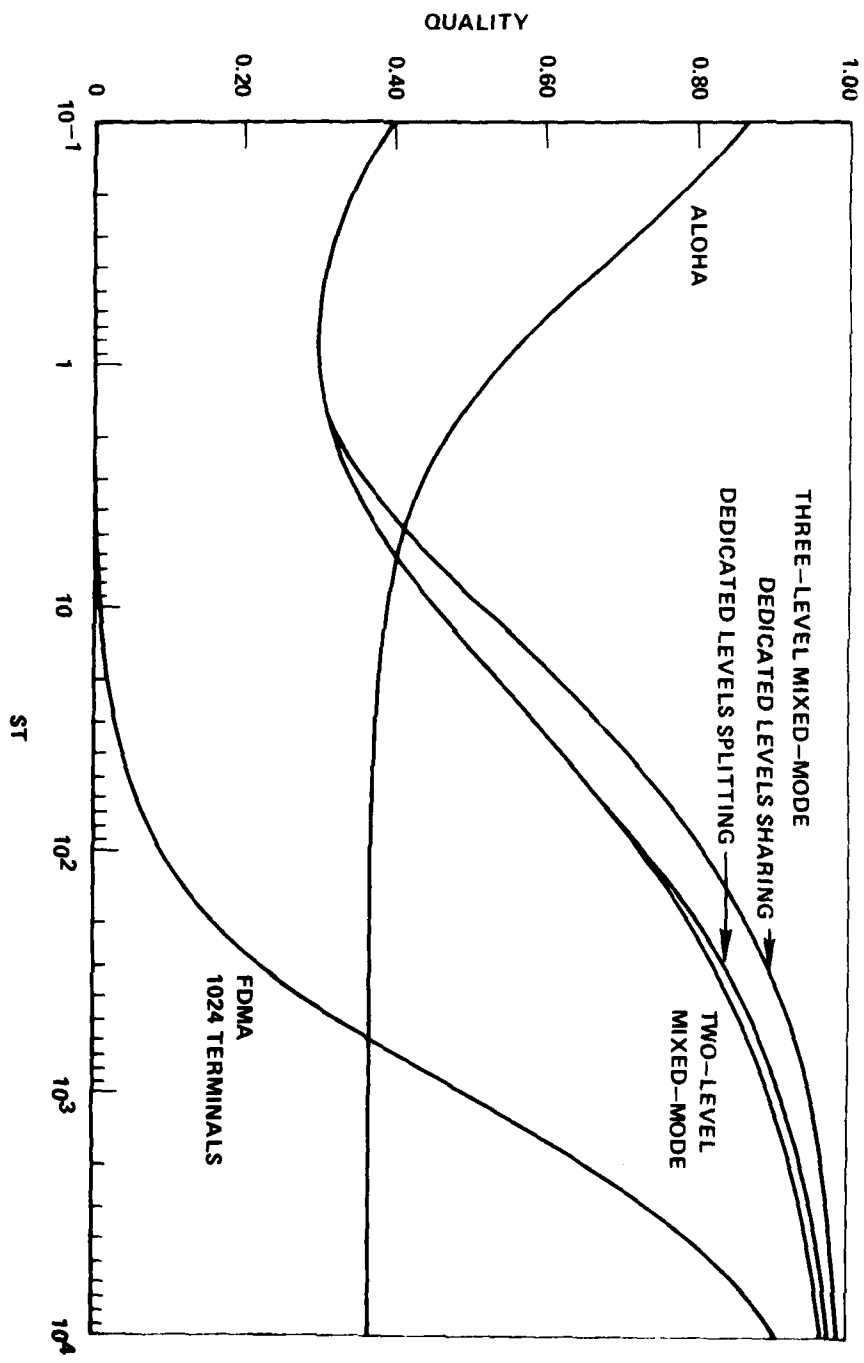


Figure 12. Quality of Three-Level Mixed-Mode Schemes, Interaction Ratio = 4.

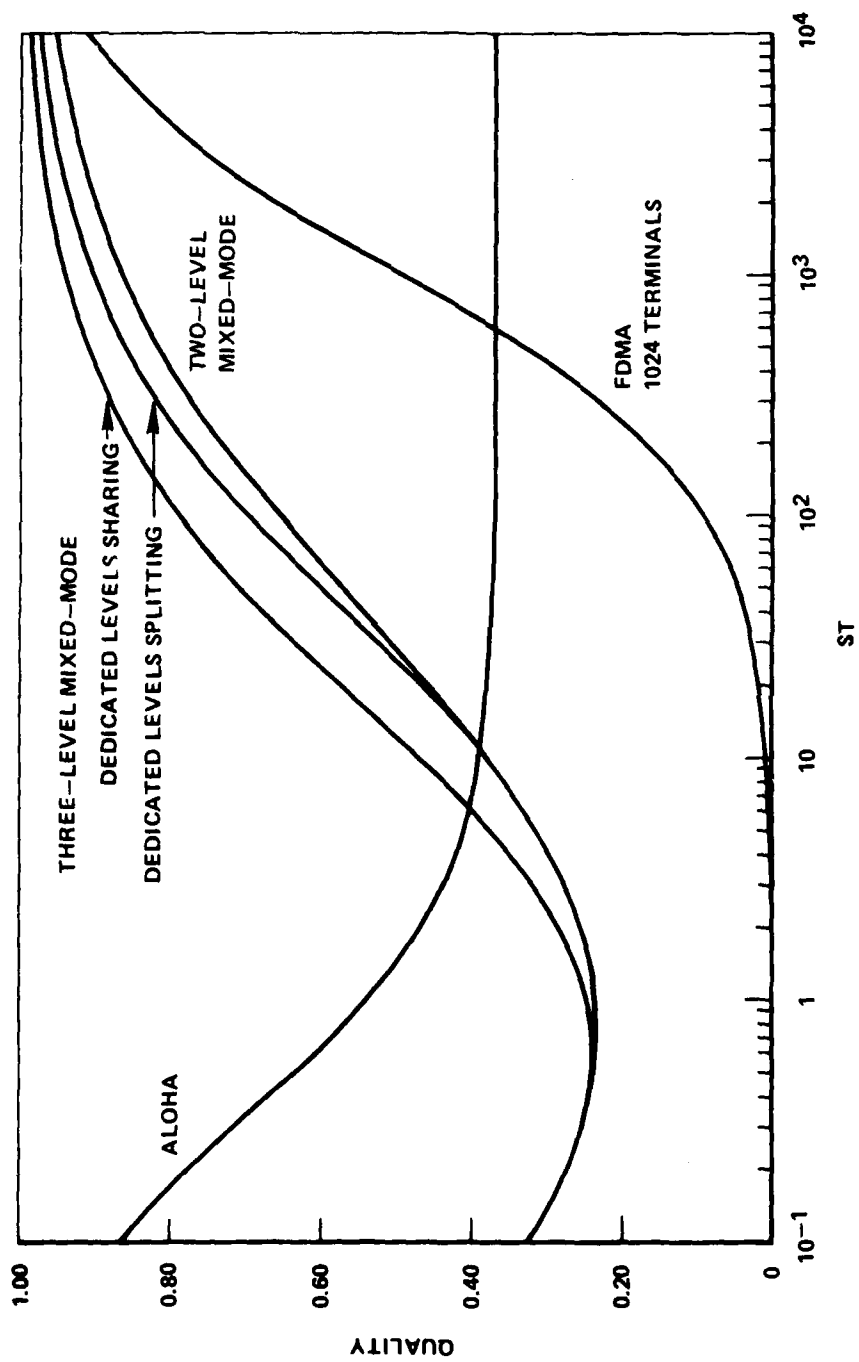


Figure 13. Quality of Three-Level Mixed-Mode Schemes, Interaction Ratio = 8.

Let us create a mixed-mode network to serve a one-dimensional system by the following procedure. Place stations at fixed intervals equal to L . Let every message go from its originating terminal to the nearest station, then over the 'station-network' to the station nearest its destination, and finally from that station to the destination itself. The connections between stations will be specified later. Dedicated broadcast channels will be used for station-station communication, and ALOHA will be used for terminal-station communication. When analyzing the mixed-mode network we shall assume that the number of terminals per station is very large, this being the worst case for random access. But when comparing with dedicated-channel networks we shall assume a set of equally spaced terminals, M occupying any section of the network with length N .

What distance is travelled by messages on the station-station level? Consider a message that has to travel a distance X from source terminal to destination terminal. The distance it will travel on the station-station (top) level depends on the location of its source terminal within the station area, but when averaging over all possible starting locations we get:

Lemma: The average station-station distance travelled by messages whose terminal-terminal distance is X , and whose starting point is uniformly distributed, is also equal to X .

Proof: X can be written as $X = nL + Y$, where n is a positive integer, L is the inter-station distance, and $0 \leq Y < L$. Let us parameterize all possible starting positions within a given station by t , where $-L/2 < t < L/2$. The distance travelled by a message on the station-station level is kL , where k is the integer nearest to $(X+t)/L$. It follows that

$$k = \begin{cases} n & \text{if } -L/2 < t < Y - L/2 \\ n+1 & \text{if } L/2 - Y < t < L/2 \end{cases}$$

The average distance travelled on the station-station level is therefore

$$n(L - Y) + (n+1)Y = nL + Y = X$$

□

It follows that the average station-station distance travelled by all messages is equal to N , the average distance between source terminal and destination terminal. This lemma does not hold for two-dimensional networks, but whenever $N \gg L$ we have that N is a good approximation to the average station-station distance travelled.

Let us assume that messages originating at one station will be heard at its nearest neighbors only, (one on each side.) What is the bandwidth necessary for such a one-level dedicated-channel network? What is a good policy for creating and allocating dedicated channels? Once we define our channels, by defining traffic streams that can be transmitted independently, the overall bandwidth necessary will depend on the capacity each channel needs and on the number of colors necessary to paint the channel so that no two of the same color interfere at their destinations.

We shall assume that every station has an omnidirectional antenna, i.e., that every transmission propagates in both directions. Two transmission policies are then possible: If all transmissions coming out of a given terminal are queued together and transmitted without regard to the direction of their destination, we need at least three colors to ensure that a terminal does not interfere with transmissions destined to itself or to its two neighbors. Three are obviously enough, because they can be assigned to terminals in a cyclic fashion. If we want transmission from a given terminal to each of its two neighbors to be done independently and at the same time, we must give each terminal two channels. Four colors are then necessary and sufficient to enable each terminal to separately send in two directions and to separately receive from two directions.

Let T_1 be the average delay suffered in the top (station-station) level by all messages, and let z be the ratio between L and N . $1/z$ is the average number of hops taken by a message in the top level. If all traffic coming out of a terminal share one channel, even though each message is destined only to one of the neighbors, the traffic on each channel is NS_u , and since three colors are necessary, we get for this case that the necessary capacity for the station-station level is

$$C_1 = 3NS_u + \frac{1}{T_1} \frac{3}{z} \quad (14)$$

If we give each terminal a separate channel for each direction, then the traffic on each channel is $NS_u/2$, and since four colors are necessary in this case we get

$$C_1 = 2NS_u + \frac{1}{T_1} \frac{4}{z} \quad (15)$$

Comparing (14) and (15), we see that it is better to have one channel per terminal when the traffic is bursty ($NS_u T \ll 1/z$) and it is better to have two channels per terminal when the traffic is steady ($NS_u T \gg 1/z$). Equations (14) and (15) will also describe one-level dedicated-channels networks if M is substituted for $1/z$, where M is the number of terminals in a portion of the network whose length is N .

Returning now to our two-level networks, we must calculate the bandwidth necessary for the bottom part. Let us first assume that all transmissions in the bottom level have a range exactly equal to $L/2$. The total traffic carried by each terminal-station system is then $2LS_u = 2zNS_u$. Despite the fact that half of this total traffic is coming from one source - the station - we shall, at first, model the bottom level by a simple ALOHA system. With our assumption on transmission range there will be no interaction between neighboring ALOHA systems, and we can write for C_2 , the capacity necessary in the bottom level,

$$C_2 = e 2zNS_u + \frac{1}{T_2} \quad (16)$$

where T_2 is the average delay for getting through a terminal-station system once.

Let us assume that separate capacities will be assigned to the terminal-station and to the station-station subsystems, without sharing. The necessary total capacity can then be obtained by minimizing $C_1 + C_2$ subject to $T_1 + 2T_2 = T$, where 2 multiplies T_2 because every message goes through two terminal-station systems, once at each end of its path.

Combining (14) and (16), for example, we get

$$C = 3NS_u + e 2zNS_u + \frac{1}{T} \left[\left(\frac{3}{z} \right) + 2 \right]^2 \quad (17)$$

Combining (15) and (16) will similarly lead to

$$C = 2NS_u + e 2zNS_u + \frac{1}{T} \left[\left(\frac{4}{z} \right) + 2 \right]^2 \quad (18)$$

The cost of the mixed-mode network can be minimized by choosing the best interstation spacing as a function of burstiness. When the traffic is bursty the best z is large, and it becomes smaller when the traffic becomes steadier.

Fig. 14 shows the quality of various one-dimensional networks. The quality of the one-dimensional ALOHA network is $(2NS_u T + 1)/(2eNS_u T + 1)$. The curve labelled 'one-level dedicated' shows the quality of the one-level organization suitable to bursty traffic (derived from (14)) when the traffic is bursty, and the one-level organization suitable to steady traffic (derived from (15)) when the traffic is steady. The two curves labelled mixed-mode bursty and mixed-mode steady were obtained.

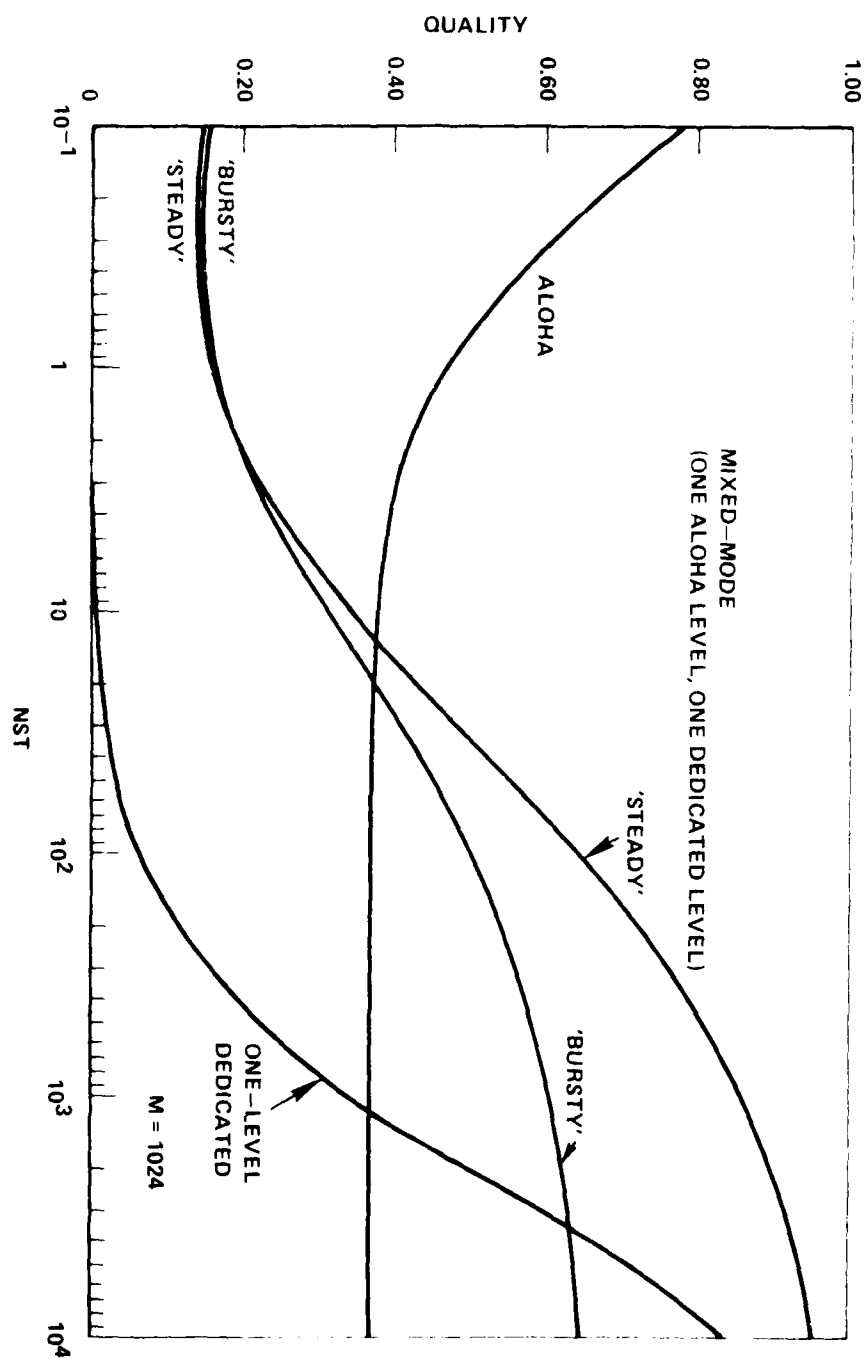


Figure 14. Mixed-Mode Networks (One ALOHA Level and One Dedicated Level).

from (17) and (18), respectively. The z that minimizes the necessary capacity was chosen for each as a function of NS_dT .

It is interesting to note that whenever the mixed-mode schemes are better than ALOHA, the steady scheme, obtained from (18), is better than the bursty one, obtained from (17). That is, the top level will be steady and should be organized accordingly.

In writing (18), we assumed that even if a message has only a small distance to go it will go to the nearest station and from that station to its destination. But the destination may be within its range, and it may be able to receive the transmission meant for the station directly. If such short-range transmissions are received directly at the end of their first hop without retransmission by a station, the system performance will be improved. Both the average number of hops necessary for messages and the amount of contending traffic in the bottom level will decrease.

It is evident a priori that this improvement will be important only when the traffic is bursty and the best interstation spacing is large. We have calculated it explicitly when the distribution of distances to be travelled by messages is exponential. Fig. 15 shows that this improvement to the mixed-mode network becomes noticeable only when the traffic is bursty enough to make the ALOHA network better than the mixed-mode network! In other words, this improvement is irrelevant.

An alternative organization for mixed-mode networks can be based on the go-forward routing policy: The first hop of each message will be to the nearest station towards its destination. The message will then use the top level to get to the last station before its destination, and then again use the bottom level to reach its destination.

If all transmissions have the same range it must be at least L in this network, and we shall assume it is exactly L . Therefore we shall have more contending traffic in the bottom level of a go-forward network than before. However, there will be less traffic using the top level, and fewer hops will be necessary there. This alternative organization will be worse than the earlier one when the traffic is steady, and will be better when the traffic is bursty. When the traffic is steady, the interstation distance will be small, and the gain in the top level will be small, but doubling the contention in the bottom level is very costly. When the traffic is bursty, contention is not a serious problem and the interstation distance is large, so the gain possible in the top level will be significant.

Fig. 15 shows the go-forward mixed-mode network (shown only when it is better than the earlier scheme); and we see that it is better only when both are worse than the ALOHA network. Organizing mixed-mode networks on the go-forward principle is never a good idea. We see here once again that when a mixed-mode network is better than ALOHA and its interstation spacing is properly chosen, its top level is 'steady'.

In the rest of this chapter when we talk about mixed-mode networks with one dedicated level we shall always refer to the mixed-mode scheme described by (18), when the best z is chosen as a function of burstiness in order to minimize the necessary capacity.

9. Improving the Random Access Part

Until now we have modelled the terminal-station level by a set of ALOHA systems. But since half the traffic in each ALOHA system is concentrated in the station it can be coordinated better than in ALOHA. What will a better terminal-station level contribute to the overall performance of the network?

AD-A157 739

PERFORMANCE TRADEOFFS AND HIERARCHICAL DESIGNS OF
DISTRIBUTED PACKET-SWIT... (U) CALIFORNIA UNIV LOS
ANGELES SCHOOL OF ENGINEERING AND APPLIED...
G AKAVIA ET AL. SEP 79 UCLA-ENG-7952

2/2

UNCLASSIFIED

F/G 17/2

ML

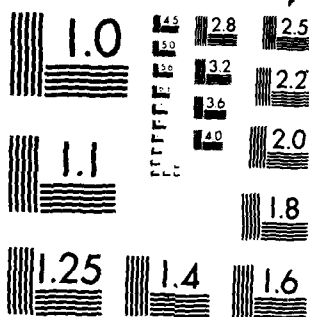
END

DATE

FILED

4-87

DTIC



MICROCOPY RESOLUTION TEST CHART
NATIONAL BUREAU OF STANDARDS 1963-A

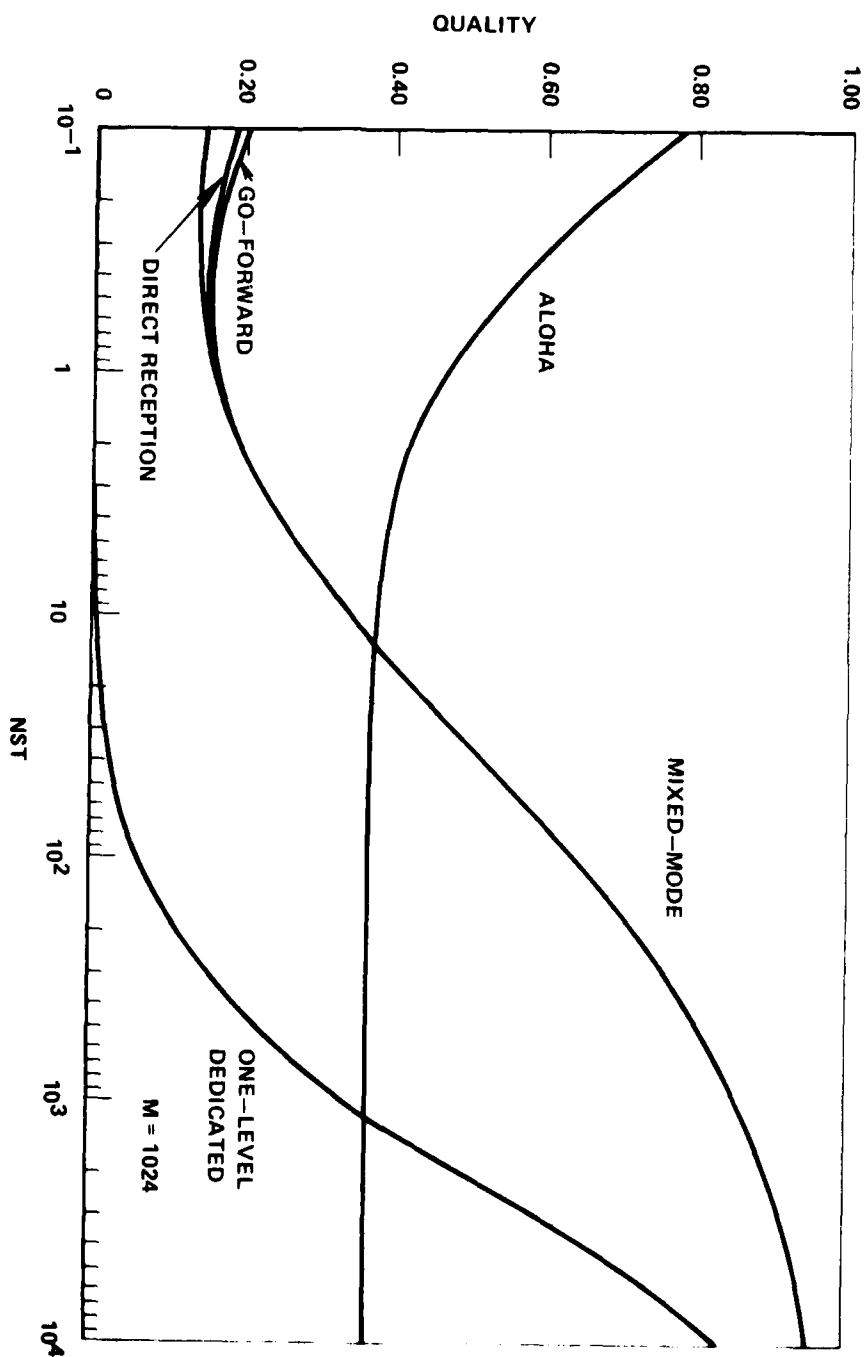


Figure 15. Alternative Mixed-Mode Networks (Direct Reception Possible, Go-Forward).

Let U be the maximum utilization of each terminal-station system. We shall model the mixed-mode network obtained with a general terminal-station access scheme by an equation similar to (18), where $1/U$ is substituted for e . We have divided the interval between $1/e$ and 1 into three equal parts, and show in Fig. 16 the quality of a mixed-mode network where U , the maximum utilization of each terminal-station system, is .367 (ALOHA), .579, .798 and 1 (M/M/1).

A mixed-mode network built with a better terminal-station access mode will obviously be better, but the improvement is not dramatic. Fig. 17 shows the ratio between the quality of a mixed-mode network with a given U and the quality of the mixed-mode network built with ALOHA as the terminal-station access mode. The curves do not go all the way to the left since they were not drawn when the mixed-mode network becomes worse than the one-level ALOHA network.

When comparing the quality of two mixed-mode networks it should be noted that the best interstation distance as a function of burstiness was chosen separately for each. This gives the mixed-mode networks an internal adjusting mechanism, and explains why improving the utilization of the terminal-station part never leads to a comparable overall improvement in the necessary capacity. When using ALOHA for the terminal-station level, we never push it to its maximum utilization, and therefore can never gain a factor of e if we assume an M/M/1 terminal-station part. We have a similar conclusion in [4] when discussing pure ALOHA networks.

Can having more than two levels improve the mixed-mode networks? By how much? We saw in [4] that a pure ALOHA network with two levels is never better than a one-level ALOHA network. But the argument used there does not apply to mixed-mode networks. A mixed-mode network with one dedicated level (the station-station level) and two ALOHA levels in the terminal-station part can lead to an improvement, but not to a large one. The maximum utilization of two-level ALOHA is .465 [4]. But even if we had a one-hop terminal-station scheme with this maximum utilization it follows from (18) that it would improve the mixed-mode network by at most 7%. Achieving this utilization by two hops will, of course, lead to an even smaller improvement.

10. More than One Dedicated Level?

More than one level in the terminal-station random access part does not lead to a significant improvement. What can we gain by having more than one level in the station-station dedicated part? What can we gain by having the optimal number of levels in the station-station dedicated part? The following is a lower bound [7] on the capacity necessary for the station-station dedicated part when the traffic is steady and when the optimal number of levels is used:

$$C_1 = 2NS_u + \frac{1}{T_1} \left[(e/2) \ln(4/z) \right]^2 \quad (19)$$

This lower bound is obtained by using regular hierarchical structures [3] to reduce the dependence of the second term of (15) on $1/z$, while ignoring the fact that when traffic is not bursty regular structures would increase the first term of (15). Combining (19) and (16) we get that the total capacity required for this mixed-mode network is

$$C = 2NS_u(1+ez) + \frac{1}{T} \left[(e/2) \ln(4/z) + 2^{1/2} \right]^2 \quad (20)$$

Fig. 18 shows the quality of a mixed-mode network with the optimal number of dedicated levels and with one dedicated level, as obtained from (20) and (18) respectively, by choosing the best z for each. Even though we use an *upper* bound on the performance of a dedicated station-station part using the *optimal* number of levels we did not gain much over the mixed-mode network that used only one dedicated level! The reason is familiar by now. Multi-level organizations are especially important when the network is *both* bursty and distributed, but this will not occur in our mixed-mode networks, since the station-station part will become very distributed (i.e., $1/z$ will become very large) only when the traffic

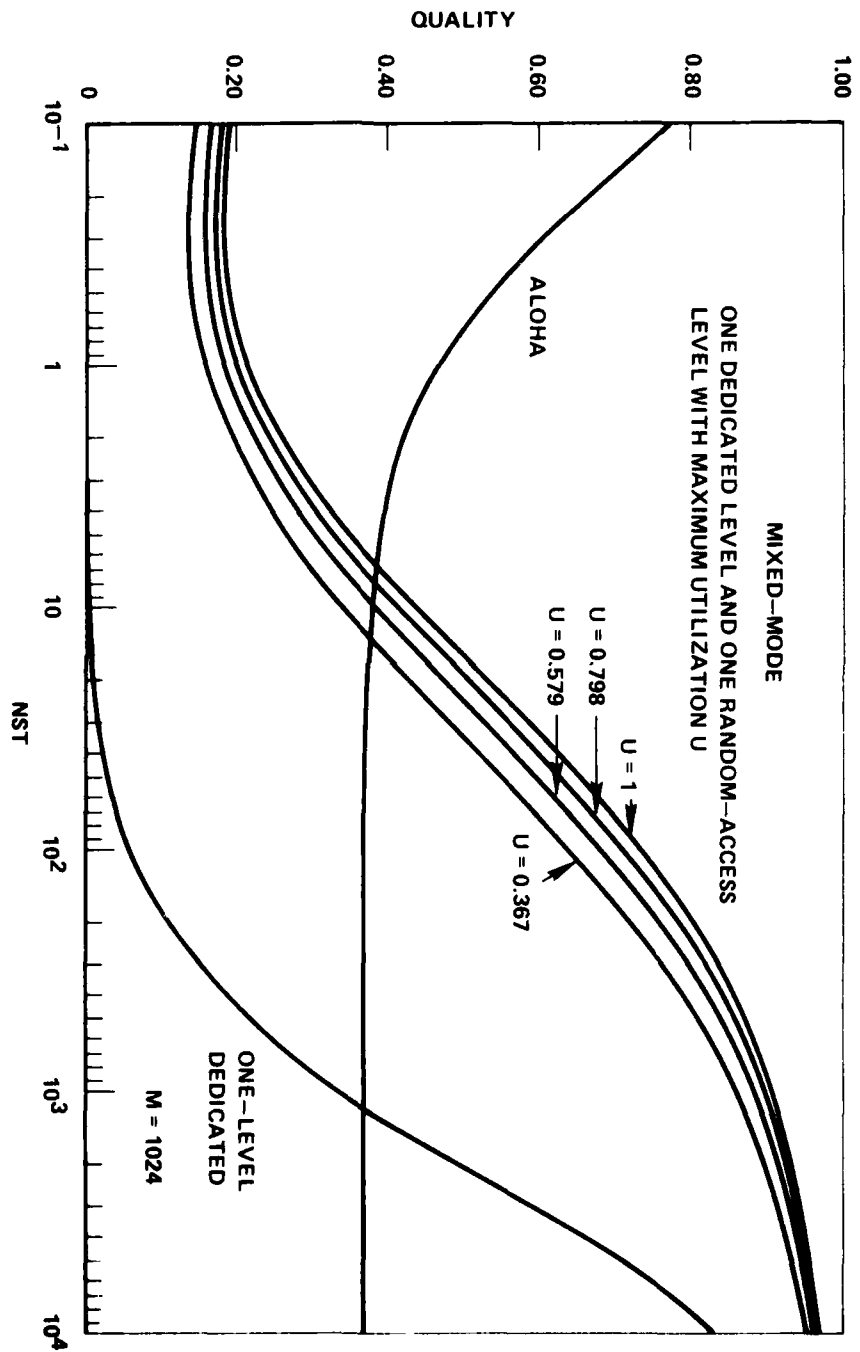


Figure 16. Mixed-Mode Network With an Improved Random Access Level.

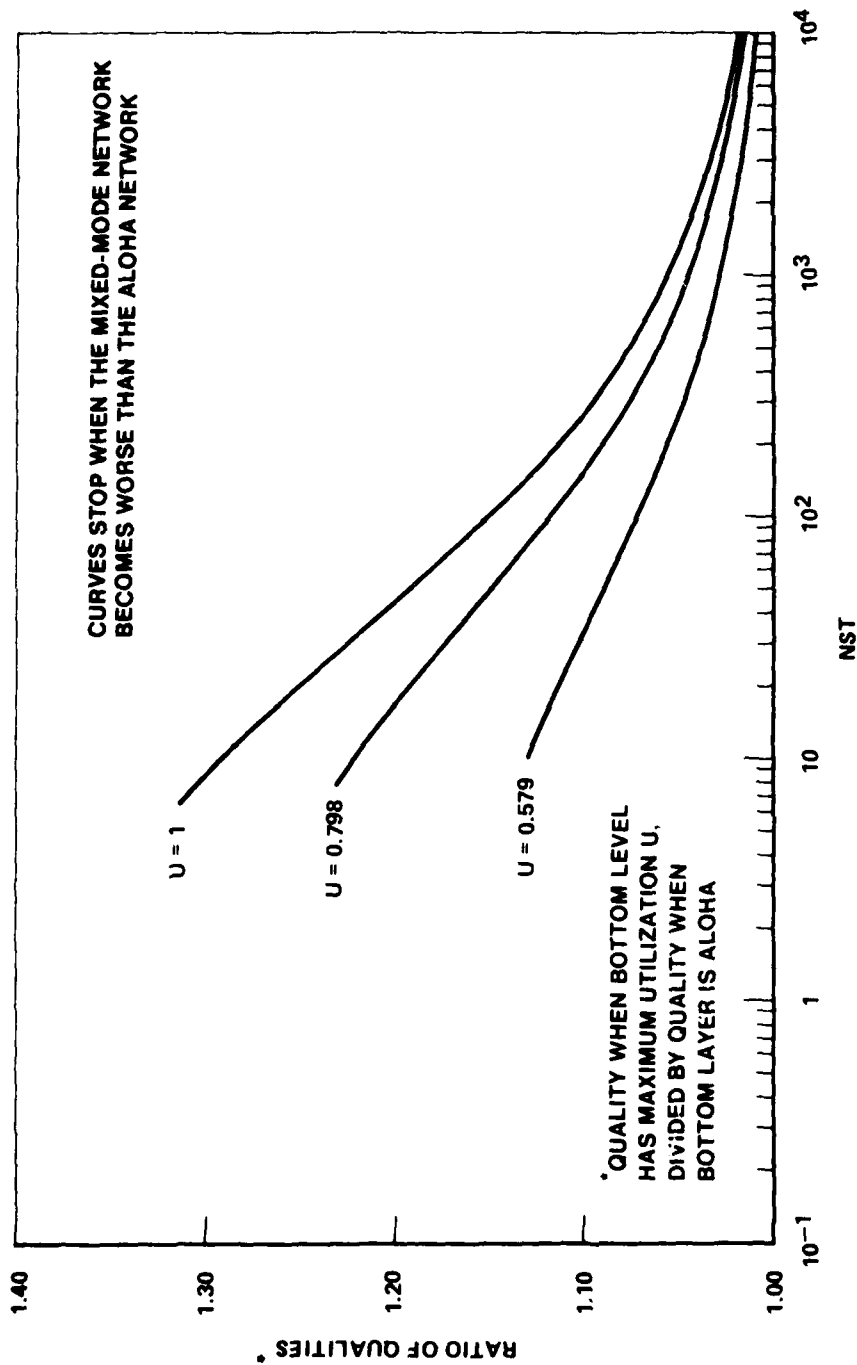


Figure 17. Gain in Overall Performance of a Mixed-Mode Network When Random Access Level is Improved.

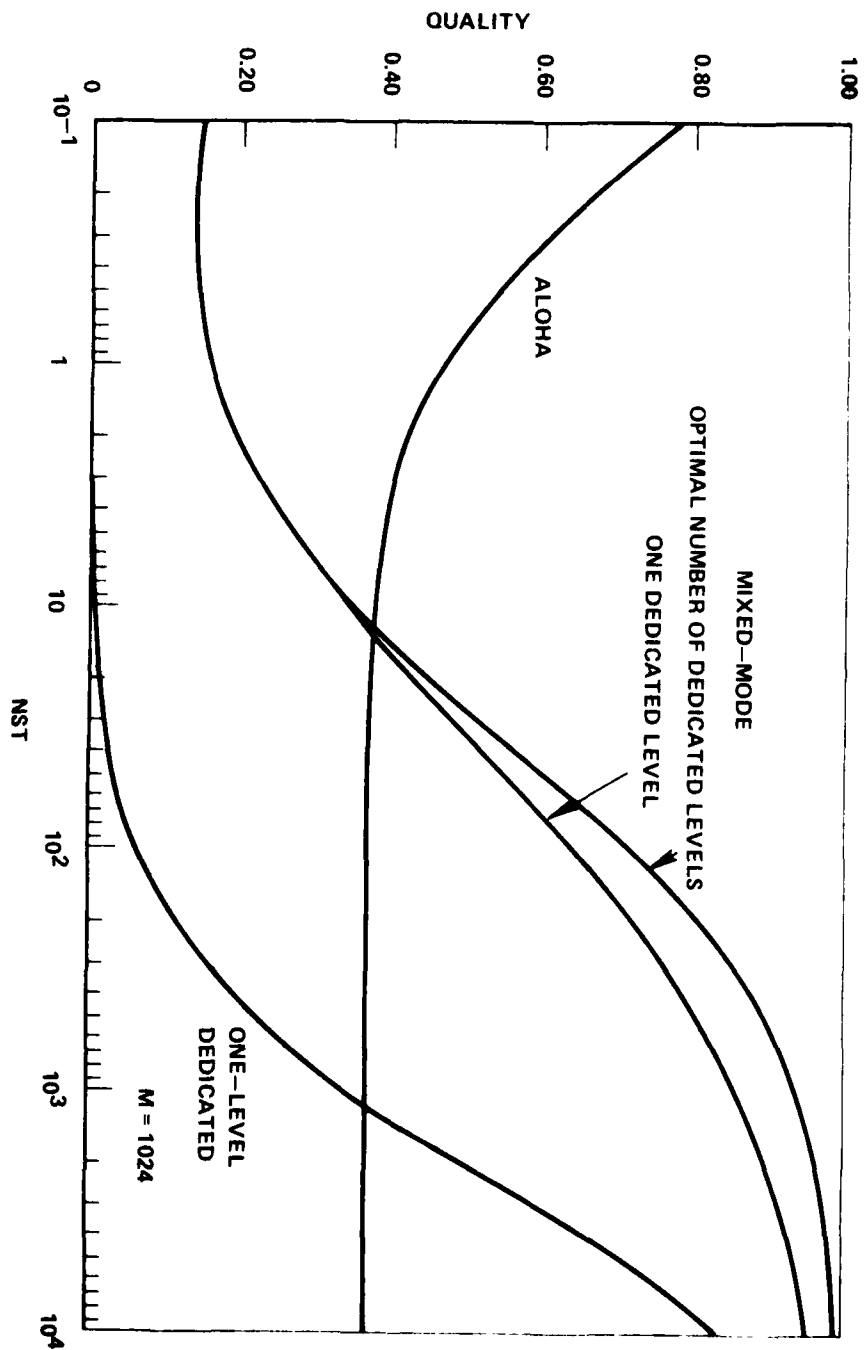


Figure 18. Mixed-Mode Network With an Optimal Number of Dedicated Levels.

is very steady.

Equation (20) can be generalized and will describe any mixed-mode network with an optimal-level dedicated station-station part and a one-level random access terminal-station part if U , the maximum utilization of the random access scheme, is substituted for $1/e$. Fig. 19 shows the ratio between the quality obtained with an optimal number of dedicated levels and the quality obtained with one dedicated level. This ratio is identically equal to 1 when the traffic is bursty because $1/z$ is then very small, and the optimal number of levels is then 1. We see from Fig. 19 that if the random access mode is better than ALOHA, introducing more than one dedicated level will lead to an even smaller improvement. Only if there is a strong interaction, and the curve with $U=.092$ can be taken to represent ALOHA with an interaction ratio equal to 4, will having more than one dedicated level lead to a more significant improvement.

11. Conclusions

ALOHA is good when the traffic is bursty, and dedicated channels are good when the traffic is steady. Mixed-mode systems, with ALOHA in the bottom level and dedicated channels in the top level, can be much better than either ALOHA or dedicated channels when the traffic is of medium burstiness and the 'amount' of mixing is properly adjusted. Under reasonably favorable conditions, the available bandwidth should be shared by the two levels, and not split between them. But even when conditions are the least favorable, and the channel must be split, the mixed-mode systems are surprisingly good.

Mixed-mode systems in general, and mixed-mode networks in particular, show a certain robustness. By the freedom to choose the right mix, the system gains an internal adjustment mechanism, and will never push any of its two parts until it is very bad. That is, the ALOHA part will never be heavily loaded and there will never be many lightly utilized dedicated channels. Because of this robustness it is harder to improve mix-mode networks. Changing the bottom level of a mixed-mode network from ALOHA to a better random access scheme leads to only a relatively small overall improvement. Introducing more dedicated levels in a mixed-mode network likewise leads to only a modest overall improvement.

References

- [1] Abramson, N., "Packet Switching with Satellites," *AFIPS Conference Proceedings*, 1973 National Computer Conference, Vol. 42, pp. 695-702.
- [2] Lam, S., "Packet Switching in a Multi-Access Broadcast Channel with Application to Satellite Communication in a Computer Network," Computer Systems Modeling and Analysis Group, School of Engineering and Applied Science, University of California, Los Angeles, UCLA-ENG-7429, April 1974. (Also published as a Ph.D. Dissertation, Computer Science Department)
- [3] Akavia, G.Y., and L. Kleinrock, "Hierarchical Use of Dedicated Channels," submitted for publication.
- [4] Kleinrock, L. and G.Y. Akavia, "On a Self Adjusting Capability of ALOHA Networks," submitted for publication.

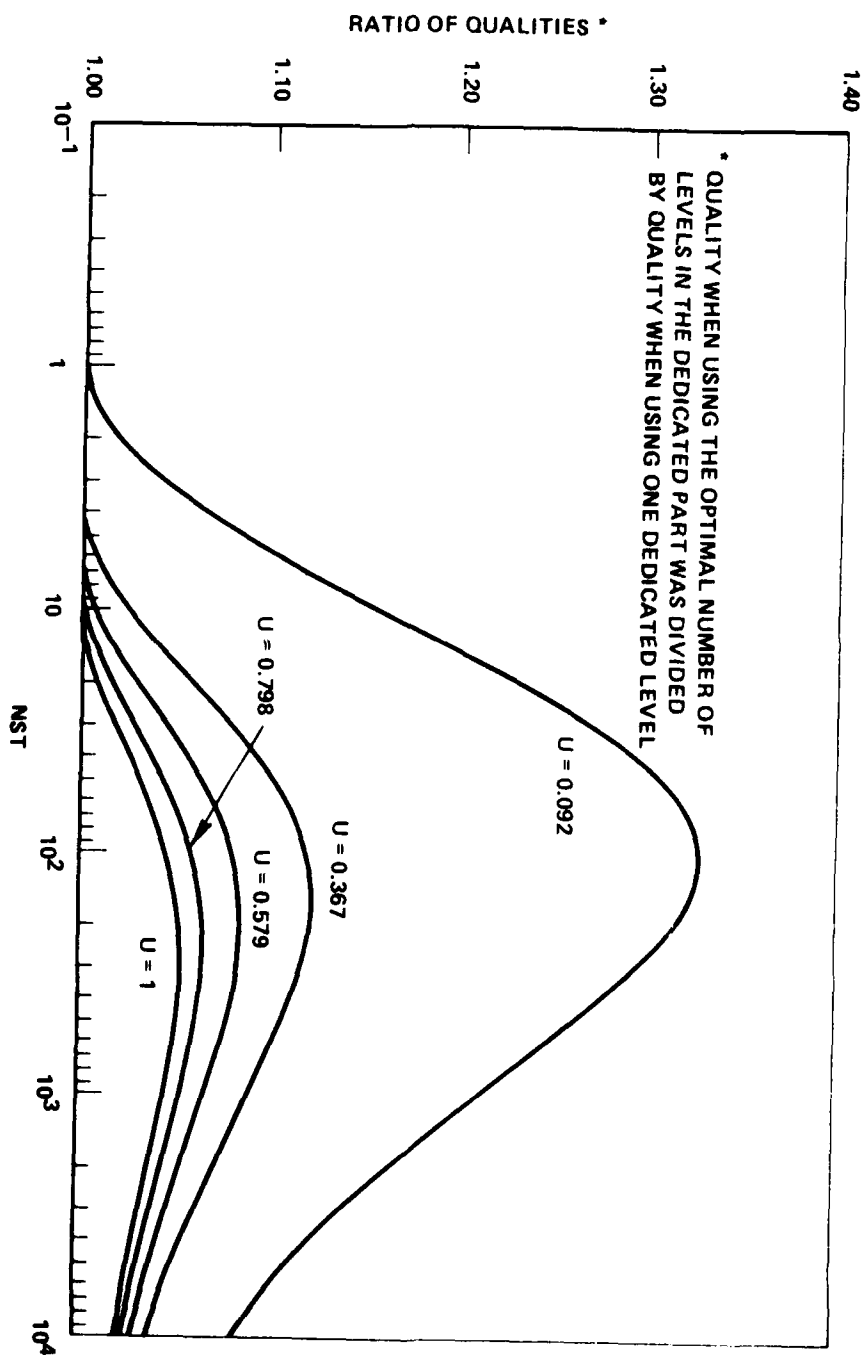


Figure 19. Gain in Overall Performance of a Mixed-Mode Network When Optimal Number of Dedicated Levels is Used.

- [5] Kleinrock, L., "Performance of Distributed Multi-Access Computer Communication Systems," *Information Processing 1977, Proceedings of IFIP Congress 77*, Toronto, August 1977, North Holland, Amsterdam, 1977, pp. 547-552.
- [6] Kleinrock, L., "Resource Allocation in Computer Systems and Computer-Communication Networks," *Information Processing 1974, Proceedings of IFIP Congress 74*, Stockholm, August 1974, North Holland, Amsterdam 1974, pp. 11-18.
- [7] Akavia, G.Y., "Hierarchical Organization of Distributed Packet-Switching Communication Systems," Ph.D. Dissertation, Computer Science Department, University of California, Los Angeles, California, 1978.
- [8] Gitman, I., "On the Capacity of Slotted ALOHA Networks and Some Design Problems," *IEEE Transactions on Communications*, Vol. COM-23, pp. 305-317, March 1975.
- [9] Roberts, L.G., "ALOHA Packets System with and without Slots and Capture," *Computer Communication Review*, A Quarterly Publication of the ACM Special Interest Group on Data Communication, Vol. 5, No. 2, pp. 28-42, April 1975.
- [10] Schiff, L., "Random-Access Digital Communication for Mobile Radio in a Cellular Environment," *IEEE Transactions on Communications*, Vol. COM-22, pp. 686-692, May 1974.

COMPUTER SYSTEMS MODELING AND ANALYSIS REPORT SERIES

Turn, R., "Assignment of Inventory of a Variable Structure Computer," January 1963, UCLA-ENG-6305.

Martin, D.F., "The Automatic Assignment and Sequencing of Computations on Parallel Processor Systems," January 1966, UCLA-ENG-6604 (AEC/ONR).

Coffman, E.G., "Stochastic Models of Multiple and Time-Shared Computer Operations," June 1966, UCLA-ENG-6638, (AEC/ARPA/ONR) NTIS AD636-976.

Bovet, D.P., "Memory Allocation in Computer Systems," June 1968, UCLA-ENG-6817 (AEC/ARPA/ONR).

Baer, J.L., "Graph Models of Computations in Computer Systems," October 1968, UCLA-ENG-6846 (AEC:UCLA-10P14-51/ ARPA/ONR) NTIS AD678-753.

Russell, E.C., "Automatic Program Analysis," March 1969, UCLA-ENG-6912 (AEC:UCLA-10P14-72/ARPA/ONR) NTIS AD 686-401.

Koster, R., "Low Level Self-Measurement in Computers," December 1969, UCLA-ENG-6957 (AEC:UCLA-10P14-84).

Cerf, V.G., "Measurement of Recursive Programs," May 1970, UCLA-ENG-7043 (AEC:UCLA-10P14-90/ARPA).

Volansky, S.A., "Graph Model Analysis and Implementation of Computational Sequences," June 1970, UCLA-ENG-7048 (AEC:UCLA-10P14-93).

Cole, G.D., "Computer Network Measurements: Techniques and Experiments," October 1971, UCLA-ENG-7165 (ARPA) NTIS AD 739-344.

Hsu, J., "Analysis of a Continuum of Processor-Sharing Models for Time-Shared Computer Systems," October 1971, UCLA-ENG-7166 (ARPA) NTIS AD739-345.

Ziegler, J.F., "Nodal Blocking in Large Networks," October 1971, UCLA-ENG-7167 (ARPA) NTIS AD 741-647.

Cerf, V.G., E. Fernandez, K. Gostelow, and S. Volansky, "Formal Control-Flow Properties of a Model of Computation," December 1971, UCLA-ENG-7178 (AEC:UCLA-10P14-105).

Gostelow, K.P., "Flow Control, Resource Allocation, and the Proper Termination of Programs," December 1971, UCLA-ENG-7223 (AEC:UCLA-10P14-106).

Cerf, V.G., "Multiprocessors, Semaphores, and a Graph Model of Computation," April 1972, UCLA-ENG-7223 (AEC:UCLA- 10P14-110).



Fultz, G.L., "Adaptive Routing Techniques for Message Switching Computer Communication Networks," July 1972, UCLA-ENG-7252 (ARPA) NTIS AD 749-678.

Fernandez, E., "Activity Transformations on Graph Models of Parallel Computations," October 1972, UCLA-ENG-7287 (AEC:UCLA-10P14-116).

Gerla, M., "The Design of Store-and-Forward (S/F) Networks for Computer Communications," January 1973, UCLA-ENG-7319 (ARPA) NTIS AD 758-704.

Tatan, R., "Optimal Control of Tandem Queues," May 1973, UCLA-ENG-7333 (AFOSR).

Postel, J.B., "A Graph Model Analysis of Computer Communications Protocols," January 1974, UCLA-ENG-7410 (ARPA) NTIS AD 777-506.

Opderbeck, H., "Measurement and Modeling of Program Behavior and its Applications," April 1974, UCLA-ENG-7418 (ONR).

Lam, S.S., "Packet Switching in a Multi-Access Broadcast Channel with Application to Satellite Communication in a Computer Network," April 1974, UCLA-ENG-7429 (ARPA) NTIS AD 781-276.

Yavne, M., "Synthesis of Properly Terminating Graphs," May 1974, UCLA-ENG-7434 (AEC:UCLA-34P214-2).

Webster, F., "An Implementation of the Burroughs D-Machine," June 1974, UCLA-ENG-7449 (AEC:34P214-6).

Sylvain, P., "Evaluating the Array Machine," August 1974, UCLA-ENG-7462 (NSF).

Kleinrock, L., "Computer Network Research Final Technical Report," August 1974, UCLA-ENG-7467 (ARPA).

Tobagi, F.A., "Random Access Techniques for Data Transmission over Packet Switched Radio Networks," December 1974, UCLA-ENG-7499 (ARPA).

White, D.E., "Fault Detection Through Parallel Processing in Boolean Algebra," March 1975, UCLA-ENG-7504 (ONR).

Wong, J.W.-N., "Queueing Network Models for Computer Systems," October 1975, UCLA-ENG-7579 (ARPA).

Loomis, M.E., "Data Base Design: Object Distribution and Resource-Constrained Task Scheduling," March 1976, UCLA-ENG-7617 (ERDA:UCLA-34P214-26).

Kamoun, F., "Design Considerations for Large Computer Communication Networks," April 1976, UCLA-ENG-7642 (ARPA).

Korff, P.B., "A Multiaccess Memory," July 1976, UCLA-ENG-7607 (NSF-OCA-MC57203633-76074).

Ng, Y.-N., "Reliability Modeling and Analysis for Fault-Tolerant Computers," September 1976, UCLA-ENG-7698 (NSF-MSC-7203633-76981).

Lee, D.D., "Simulation Study of a Distributed Processor Computer Architecture," November 1976, UCLA-ENG-76106 (ONR).

Scholl, M.O., "Multiplexing Techniques for Data Transmission over Packet Switched Radio Systems," December 1976, UCLA-ENG-76123 (ARPA).

Erlach, Z., "On Centralized Bus Transportation Systems with Poisson Arrivals," December 1976, UCLA-ENG-76124 (ARPA).

Naylor, W.E., "Stream Traffic Communication in Packet Switched Networks," August 1977, UCLA-ENG-7760 (ARPA).

Kermani, P., "Switching and Flow Control Techniques in Computer Communication Networks," February 1978, UCLA-ENG-7802 (ARPA).

Akavia, G. and Kleinrock, L., "Performance Tradeoffs and Hierarchical Designs of Distributed Packet Switching," September 1979, UCLA-ENG-7952 (ARPA).

DAT
ILMI