

FILE COPY

COMPUTER SYSTEMS LABORATORY

STANFORD UNIVERSITY · STANFORD, CA 94305-2192

AD-A221 791

87-0246

CHANNEL ACCESS SCHEMES AND
FIBER OPTIC CONFIGURATIONS
FOR INTEGRATED-SERVICES
LOCAL AREA NETWORKS

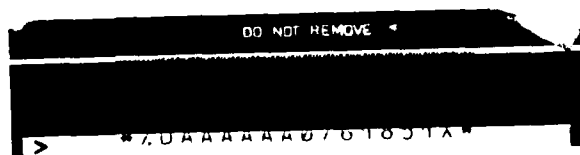
M. Mehdi Nassehi

Technical Report: CSL-TR-87-322

March 1987



This report is the author's Ph.D. dissertation which was completed under the advisorship of Professor Fouad A. Tobagi. This work was supported by the National Aeronautics and Space Administration under Grants No. NAG2-292 and No. NAGW-419, the Defense Advanced Research Projects Agency under Contract No. MDA 903-79-C-0201 and an IBM Graduate Fellowship.



90 05 14 125

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

REPORT DOCUMENTATION PAGE		READ INSTRUCTIONS BEFORE COMPLETING FORM
1. REPORT NUMBER 87-322	2. GOVT ACCESSION NO.	3. RECIPIENT'S CATALOG NUMBER
4. TITLE (and Subtitle) CHANNEL ACCESS SCHEMES AND FIBER OPTIC CONFIGURATIONS FOR INTEGRATED-SERVICES LOCAL AREA NETWORKS		5. TYPE OF REPORT & PERIOD COVERED TECHNICAL REPORT
		6. PERFORMING ORG. REPORT NUMBER 87-322
7. AUTHOR(s) M. Mehdi Nassehi		8. CONTRACT OR GRANT NUMBER(s) NASA NAG2-292 NASA NAGW-419
9. PERFORMING ORGANIZATION NAME AND ADDRESS Stanford Electronics Laboratory Stanford University Stanford, CA 94305-2192		10. PROGRAM ELEMENT, PROJECT, TASK AREA & WORK UNIT NUMBERS
11. CONTROLLING OFFICE NAME AND ADDRESS National Aeronautics and Space Administration Washington, D.C. 20546		12. REPORT DATE March 1987
		13. NUMBER OF PAGES 161
14. MONITORING AGENCY NAME & ADDRESS (if different from Controlling Office) Resident Representative Office of Naval Research Durand 165 Stanford University, Stanford, CA 94305-2192		15. SECURITY CLASS. (of this report) UNCLASSIFIED
		15a. DECLASSIFICATION/DOWNGRADING SCHEDULE
16. DISTRIBUTION STATEMENT (of this Report) Approved for public release; unlimited distribution.		
17. DISTRIBUTION STATEMENT (of the abstract entered in Block 20, if different from Report)		
18. SUPPLEMENTARY NOTES		
19. KEY WORDS (Continue on reverse side if necessary and identify by block number)		
20. ABSTRACT (Continue on reverse side if necessary and identify by block number) Local Area Networks have been in common use for data communications for several years and have enjoyed great success. More recently, there has been a growing interest in using a single network to support many applications (e.g., speech, high-resolution graphics, facsimile, video, etc.) in addition to traditional data traffic. This leads to so-called Integrated Services Local Area Networks. These additional applications introduce new requirements in terms of volume of traffic and real-time delivery of data which are not met by exist-		

DD FORM 1 JAN 73 1473

EDITION OF 1 NOV 65 IS OBSOLETE
S N 0102-LF-014-6601

UNCLASSIFIED

SECURITY CLASSIFICATION OF THIS PAGE (When Data Entered)

ing networks. To satisfy these requirements, one needs a high-bandwidth transmission medium, such as fiber optics, and a distributed channel access scheme for the efficient sharing of the bandwidth among the various applications.

As far as the throughput-delay requirements of the various applications are concerned, a network structure along with a distributed channel access scheme are proposed which incorporate appropriate scheduling policies for the transmission of outstanding messages on the network. The proposed solution is developed in two steps. First, considering each message to be assigned a delay-cost function based on the application to which it belongs, and assuming perfect knowledge of all outstanding messages, a dynamic scheduling policy is devised which outperforms all existing policies in terms of minimizing the expected cost per message. Secondly, as required for the distributed implementation of the scheduling policy, a broadcast mechanism is devised for the efficient dissemination of all relevant information. The broadcast mechanism is based on unidirectional network structures which provide ordering among the stations.

As far as the high-bandwidth transmission medium is concerned, fiber optic technology is considered in this work. The physical ordering among stations which is required by the above access scheme may be achieved either by an active configuration, such as a ring, or by a passive configuration, such as unidirectional bus. In rings, the use of fiber optics does not introduce any new technical problems, since there exist only point-to-point links between neighboring stations. However, in the multi-tapped linear bus, the reciprocity and excess loss of optical couplers along with the low impedance of optical detectors severely limit the number of stations that can be accommodated. A number of alternative passive configurations are proposed which overcome this limitation. Also presented is a unified method for selecting coupler coefficients to *minimize the maximum path loss*. This method is used to compute the maximum number of stations that a configuration can support.

CHANNEL ACCESS SCHEMES AND FIBER OPTIC CONFIGURATIONS FOR INTEGRATED-SERVICES LOCAL AREA NETWORKS

M. Mehdi Nassehi

Technical Report: CSL-TR-87-322

March 1987

Computer Systems Laboratory
Departments of Electrical Engineering and Computer Science
Stanford University
Stanford, California 94305-2191

Accession For	
NTIS	GRANT
DTIC	TR
Unannounced	
Justification	
By	
Disc	
A	
Dist	
A-1	

Abstract

Local Area Networks have been in common use for data communications for several years and have enjoyed great success. More recently, there has been a growing interest in using a single network to support many applications (e.g., speech, high-resolution graphics, facsimile, video, etc.) in addition to traditional data traffic. This leads to so-called Integrated Services Local Area Networks. These additional applications introduce new requirements in terms of volume of traffic and real-time delivery of data which are not met by existing networks. To satisfy these requirements, one needs a high-bandwidth transmission medium, such as fiber optics, and a distributed channel access scheme for the efficient sharing of the bandwidth among the various applications.

As far as the throughput-delay requirements of the various applications are concerned, a network structure along with a distributed channel access scheme are proposed which incorporate appropriate scheduling policies for the transmission of outstanding messages on the network. The proposed solution is developed in two steps. First, considering each message to be assigned a delay-cost function based on the application to which it belongs, and assuming perfect knowledge of all outstanding messages, a dynamic scheduling policy is devised which outperforms all existing policies in terms of minimizing the expected cost per message. Secondly, as required for the distributed implementation of the scheduling policy, a broadcast mechanism is devised for the efficient dissemination of all relevant information. The broadcast mechanism is based on unidirectional network structures which provide ordering among the stations.

As far as the high-bandwidth transmission medium is concerned, fiber optic technology is considered in this work. The physical ordering among stations which is required by the above access scheme may be achieved either by an active configuration, such as a ring, or by a passive configuration, such as unidirectional bus. In rings, the use of fiber optics does not introduce any new technical problems, since there exist only point-to-point links between neighboring stations. However, in the multi-tapped linear bus, the reciprocity and excess loss of optical couplers along with the low impedance of optical detectors severely

limit the number of stations that can be accommodated. A number of alternative passive configurations are proposed which overcome this limitation. Also presented is a unified method for selecting coupler coefficients to minimize the maximum path loss. This method is used to compute the maximum number of stations that a configuration can support.

Copyright ©1987 M. Mehdi Nassehi

Acknowledgments

I wish to express my sincere appreciation to my advisor, Professor Fouad Tobagi, for his invaluable guidance in all aspects of this thesis; his advice offered on numerous occasions is also greatly appreciated. I am grateful to Professor Michel Marhic of Northwestern University for his guidance in the work done on fiber optics.

I gratefully acknowledge IBM's Watson Research Center for supporting me with a fellowship for 2 years. My research was also supported by NASA through a research contract under the direction of Professor Tobagi.

To Jill Sigl, my appreciation for assisting me in preparing this thesis. I would like to thank my friends, Yitzhak Birk, José Brázio, Ali Grami, Masood Namjoo, and David Shur, for helping me in the difficult parts of this work. I am grateful to all members of my family for their moral support. In particular, I would like to express my gratitude to my parents, for their love and encouragement.

Contents

1	Introduction	1
1.1	Local Area Networks	1
1.2	Integrated-Services	6
1.3	Contributions	7
2	Transmission Scheduling Policies for Broadcast Local Area Networks	12
2.1	The Transmission Scheduling Problem	13
2.2	Scheduling Policies	16
2.2.1	Basic Scheduling Policies	16
2.2.2	A Policy based on Static Scheduling	17
2.2.3	Dynamic-Priority Scheduling Policy	23
2.3	Performance	27
2.3.1	Traffic Model	30
2.3.2	Simulation	31
2.3.3	Numerical Results	32
2.4	Summary	48
3	Broadcast Local Area Networks Incorporating Scheduling Policies	49
3.1	Attempt-and-Defer Access Schemes	50

3.2	Basic Requirements for the Distributed Implementation of Scheduling Policies	53
3.3	A Design Meeting The Requirements	55
3.3.1	Broadcast Mechanism	55
3.3.2	Station Architecture	57
3.4	Summary	60
4	Fiber Optic Configurations for Local Area Networks	61
4.1	Feature Abstraction	62
4.2	Fiber Optic Configurations	65
4.2.1	Configurations for the Control Subnetwork	65
4.2.1.1	Stretched Configuration	65
4.2.1.2	Bypassed Configuration	66
4.2.1.3	Non-reciprocal Couplers	67
4.2.2	Configurations for the Data Subnetworks	68
4.2.2.1	Basic Configurations	68
4.2.2.2	Hybrid Configurations	69
4.2.2.3	Compound Configurations	71
4.3	Power Budget Analysis	72
4.3.1	Formulation of the Problem	72
4.3.2	Optimization of Couplers	74
4.3.3	A Bound on Number of Stations	83
4.3.4	Coupler Optimization in linear subnetworks	87
4.4	Performance	90
4.4.1	Configurations for Control Subnetwork	91
4.4.2	Configurations for Data Subnetwork	94
4.5	Summary	97

5	Conclusions and Future Research	99
5.1	Conclusions	99
5.2	Suggestions for Additional Research	102
A	Additional Results on Fiber Optic Configurations	105
A.1	Fiber-Optics Components	105
A.2	Description of Configurations	109
A.3	Uniform Optimization of Couplers	110
A.4	Numerical Results	119
A.4.1	Commercially Available Components	119
A.4.2	Maximum Number of Stations	121
A.4.3	Comments on the Optimum Coupling Fractions	126
A.4.4	Sensitivity of $N_{\max}$ to Component Parameters	128
	References	140

Figures

Fig. 2.1	Examples of delay cost function for the n th message generated in the system.	15
Fig. 2.2	Delay cost and urgency versus the earliest possible completion time for the step cost function.	28
Fig. 2.3	Delay cost and urgency versus the earliest possible completion time for the ramp cost function.	29
Fig. 2.4	Expected cost per message versus normalized discounting factor, for different combinations of message-length distributions and cost functions.	33
Fig. 2.5	Expected cost per message versus normalized mean delay allowance, for fixed-length messages and step cost functions.	35
Fig. 2.6	Expected cost per message versus normalized mean delay allowance, for fixed-length messages and ramp cost functions.	36
Fig. 2.7	Expected cost per message versus normalized mean delay allowance, for exponentially-distributed message lengths and step cost functions.	37
Fig. 2.8	Expected cost per message versus normalized mean delay allowance, for exponentially-distributed message lengths and ramp cost functions.	38
Fig. 2.9	Expected cost per message versus load factor ρ for fixed-length messages and step cost functions.	40
Fig. 2.10	Expected cost per message versus load factor ρ for fixed-length messages and ramp cost functions.	41

Fig. 2.11	Expected cost per message versus load factor ρ for exponentially-distributed message lengths and step cost functions.	42
Fig. 2.12	Expected cost per message versus load factor ρ for exponentially-distributed message lengths and ramp cost functions.	43
Fig. 2.13	Expected cost per message versus normalized maximum packet length (k/L) ρ for fixed-length messages and step cost functions.	44
Fig. 2.14	Expected cost per message versus normalized maximum packet length (K/L) for fixed-length messages and ramp cost functions.	45
Fig. 2.15	Expected cost per message versus normalized maximum packet length (K/L) for exponentially-distributed message lengths and step cost functions.	46
Fig. 2.16	Expected cost per message versus normalized maximum packet length (K/L) for exponentially-distributed messages lengths and ramp cost functions.	47
Fig. 3.1	Unidirectional bus structures.	51
Fig. 3.2	Slot Format.	56
Fig. 3.3	Station architecture	58
Fig. 4.1	Star configuration.	62
Fig. 4.2	Configurations for the control subnetwork.	65
Fig. 4.3	Optical coupler.	67
Fig. 4.4	Basic configurations for the data subnetwork.	70
Fig. 4.5	Hybrid configurations for the data subnetwork.	71
Fig. 4.6	Compound configurations for the data subnetwork.	73
Fig. 4.7	The partitioning of a configuration by a coupler C_i and the corresponding notation.	76
Fig. 4.8	$\underline{G}(S_i; \mathbf{x})$ as a function of x_i .	77
Fig. 4.9	A configuration in which local optimality of the couplers is not necessary for global optimality.	79

Fig. 4.10	A loop-free configuration which does not satisfy condition C2.	80
Fig. 4.11	A collection subnetwork based on a (general) binary tree topology.	80
Fig. 4.12	Linear collection subnetwork.	82
Fig. 4.13	A configuration which does not satisfy condition C1.	82
Fig. 4.14	A general compound data subnetwork.	84
Fig. 4.15	The maximum number of stations as a function of the funneling width in a general data subnetwork. It is assumed that the components are ideal and the couplers are individually optimized.	86
Fig. 4.16	The trade-off between the maximum number of stations and fiber length in TST and LSL data subnetworks. It is assumed that the components are ideal and the couplers are individually optimized.	88
Fig. 4.17	Maximum number of stations versus power margin in the control subnetwork for a coupler excess loss of 1dB.	92
Fig. 4.18	Maximum number of stations versus coupler excess loss in the control subnetwork for a power margin of 40 dB.	93
Fig. 4.19	Maximum number of stations versus power margin in the basic and hybrid configurations for a coupler excess loss of 1 dB.	95
Fig. 4.20	Trade-off between maximum number of stations and fiber length in the compound configurations for a coupler excess loss of 1 dB and a power margin of 40 dB.	98
Fig. A.1	Fiber-optics components: (a) a connector, (b) a joint, (c) a coupler, and (d) an $S \times S$ star coupler.	107
Fig. A.2	Linear configuration (L).	109
Fig. A.3	Star configuration (S/C).	110
Fig. A.4	Tree configuration (T/C).	110
Fig. A.5	Linear/star configuration (LS).	111
Fig. A.6	Linear/tree configuration (LT).	111
Fig. A.7	Linear/star/linear configuration (LSL/C).	112

Fig. A.8	Tree/star/tree configuration (TST/C).	112
Fig. A.9	The subnetworks on which the uniform optimization is performed.	114
Fig. A.10	$\underline{G}(S; x)$ as a function of x for linear control subnetwork and $N = 3$.	115
Fig. A.11	$\underline{G}(S; x)$ as a function of x for linear control subnetwork and $N \geq 4$.	116
Fig. A.12	$\underline{G}(S; x)$ as a function of x for linear data collection subnetwork.	119
Fig. A.13	The maximum number of stations as a function of the funneling width for multimode fiber and individually optimized couplers.	122
Fig. A.14	The trade-off between the maximum number of stations and fiber length for multimode fiber and individually optimized couplers.	123
Fig. A.15	Individually optimized coupling fractions in the linear subnetworks for $N = 20$, multimode fiber, and MR components.	127
Fig. A.16	$N_{\max}$ in the T subnetwork as function of the uniform coupling fraction.	129
Fig. A.17	$N_{\max}$ in the C subnetwork as a function of the uniform coupling fraction.	130
Fig. A.18	The maximum number of stations as a function of the power margin for HC components, multimode fiber, and individually optimized couplers.	131
Fig. A.19	The maximum number of stations as a function of the power margin for MR components, multimode fiber, and individually optimized couplers.	132
Fig. A.20	The maximum number of stations as a function of the power margin for LC components, multimode fiber, and individually optimized couplers.	133
Fig. A.21	The maximum number of stations as a function of the power margin for HC components, multimode fiber, and uniformly optimized couplers.	135

Fig. A.22	The maximum number of stations as a function of the coupler excess-loss for multimode fiber and individually optimized couplers. In the S subnetwork MR components are assumed while in the other subnetworks HC components are assumed.	137
Fig. A.23	The maximum number of stations as a function of the joint insertion loss for multimode fiber and individually optimized couplers. In the S subnetwork MR components are assumed while in the other subnetworks HC components are assumed.	138
Fig. A.24	The maximum number of stations as a function of the fiber attenuation for multimode fiber and individually optimized couplers. In the S subnetwork MR components are assumed while in the other subnetworks HC components are assumed.	139

Tables

Table 2.1	Example illustrating the operation of algorithm 2.1.	20
Table 2.2	Example illustrating the operation of Algorithm 2.2.	22
Table A.1	The excess-loss for a family of $S \times S$ star couplers.	108
Table A.2	Values used for fiber-optics component parameters.	120
Table A.3	Maximum number of stations for a data rate of 50 Mbps.	124
Table A.4	Maximum number of stations for a data rate of 100 Mbps.	125
Table A.5	Maximum number of stations for a data rate of 200 Mbps.	125
Table A.6	$N_{\max}$ in the LT configuration.	128

Chapter 1

Introduction

1.1 Local Area Networks

The concept of packet switching was first introduced in 1964 by Baran [BARA64] as an alternative to circuit switching in military communications. Before the end of the 60's packet switching became the primary switching method in computer communications. This technique is useful when the traffic is bursty, i.e., when the peak data rate to average data rate is high. In packet-switched networks, rather than dedicating a circuit to a pair of users for the duration of a session, messages are transmitted in separate packets. The communication lines are utilized by a pair of users only for the duration of each packet transmission. Each packet may traverse the path from the source to the destination in several hops, with storage at intermediate nodes in *store-and-forward* networks.

Now that we are entering the third decade of computer communications, packet switching has become widely used and perfected. The earliest network that utilized this packet switching technology is a terminal access network designed and implemented at the National Physics Laboratory (NPL) in England [DAVI67]. But the most prominent packet switching network is the Arpanet [KLEI76], the design and implementation of which began in 1969. Although the NPL packet switched network was used for the connection of terminals and computers and was thus a local area network, it was not until early 70's that the field of local area networks became known as such and became one of the fastest growing areas in computer communications. Local area networks are intended to bring the advantages of packet switching communication technology to bear on the communications needs of the expanding in-house computer facilities. They can generally be differentiated from long-haul store-and-forward networks by the following characteristics: local area networks span a geographical area of a few kilometers such as a single building or cluster of buildings; the data rate on a local area network is high, typically 1-100Mb/s; such networks are usually privately owned rather than publicly available; the cost of transmission and controlling the network is low compared with long-haul networks.

Existing local area networks can be broadly categorized into two basic types. These are broadcast busses and ring systems [TANE81, CLAR78, STAL84, KUMMS2]. In ring systems the data flow is unidirectional, propagating around the ring from station to station. The interface between a station and the network is an active device which receives the signal from the incoming line and retransmits it on the outgoing line. Various techniques for accessing the channel exist which give rise to various types of ring networks such as token rings, slotted rings, and register insertion rings.

Multiaccess broadcast bus systems have been popular due to the fact that by

combining the advantages of packet switching with those of broadcast communications, they offer efficient solutions to the communication needs of the local area environment both in simplicity of topology and flexibility in satisfying growth. Random access methods such as Carrier Sense Multiple Access (CSMA) [TOBA80, KLEI75] have been effectively employed. The Ethernet [METC76] is a common example. In these systems, a station that wishes to transmit senses the channel to see if it is currently being used. If the channel is busy, the station in question reschedules its transmission for some later time. If the channel is idle, the packet is transmitted. Note that it takes a finite amount of time for the signal to propagate across the network. Therefore, it is possible for a station to detect the channel to be idle and to begin transmission when, in fact, another station has already begun to transmit. In this circumstance, two transmissions occur simultaneously and a *collision* occurs. The efficiency of the channel can be improved if stations listen for such collisions and abort their transmissions when a collision occurs [TOBA79]. This variant of the access protocol is called CSMA with collision detection (CSMA/CD).

Analysis of ring and CSMA networks have shown that their performance is function of the ratio

$$a = \frac{\tau W}{B}$$

where τ is the propagation delay (in seconds) of the signal around the ring, or the end-to-end propagation delay in a CSMA bus, W is the bandwidth of the channel in bits per second, and B is the number of bits in a packet. In particular, the network capacity (i.e., maximum achievable throughput) decreases significantly as a increases.

Recently, a number of new schemes have been proposed for broadcast bus networks. These schemes provide conflict-free transmission using distributed access protocols with *round robin scheduling functions* which thus lead to bounded delay.

The stations that are "alive" are ordered so as to form what is called a *logical ring* according to which they are given their chance to transmit. We refer to these network proposals as demand assignment multiple access (DAMA) schemes. In some of these schemes, such as the Token-Passing Bus Access Method [IEEE83], an *explicit message* gets sent around the logical ring to provide the required scheduling; the station holding the token at any instant is the one that has access to the channel at that instant. It relinquishes its right to access the channel by transmitting the token to the next one in turn. Unfortunately, here too the performance degrades significantly with a [STAL84].

In contrast to those schemes where a station transmits an explicit token to the next in turn, in others the stations rely on various events due to activity on the channel to determine when to transmit. Since the token passing operation is *implicit*, the overall robustness of the network is improved over token bus networks. Here too, packet delay is bounded; but in addition both throughput and delay can be made much less sensitive to a , thus rendering these schemes particularly suitable to networks with high bandwidth, small size packets (such as those arising from real time applications), and long distances. A comprehensive survey of implicit token DAMA schemes can be found in [FINE84].

Among the various DAMA schemes, a particularly attractive class has been identified and is known as the attempt-and-defer class of access protocols. These employ a unidirectional bus structure in which signals propagate in only one direction. (Incidentally, fiber optics technology leads to unidirectional busses.)

Broadcast communications is achieved in two ways. The first approach is to fold a unidirectional cable onto itself so as to create two channels, an *outbound channel* onto which the users transmit packets and an *inbound channel* from which users receive packets; all signals transmitted on the outbound channel are returned

on the inbound channel. This is referred to as the *folded bus* structure. Another way to achieve broadcast communications is to provide two unidirectional buses with signals propagating in opposite directions. This is referred to as the *dual-bus* structure.

In both the folded-bus and dual-bus structures, on every unidirectional bus segment to which stations are connected, there exists an implicit ordering among the stations which the DAMA schemes under consideration make use of. These schemes also require that on every unidirectional bus segment, each station be given the ability to sense activity due to stations on the upstream side of its transmit tap. In the dual-bus structure, the receive taps provide this function while in the folded-bus structure additional sense taps are needed.

The attempt-and-defer access mechanism basically operates as follows. A station wishing to transmit on a given bus waits until that bus is idle. It then begins to transmit, thus establishing its desire to acquire the bus. However, if another transmission from upstream is detected, the station aborts its transmission and defers to the one from upstream. The most upstream transmission is therefore allowed to continue conflict free. Examples of networks using the attempt-and-defer access mechanism are Expressnet [TOBA83], Fasnet [LIMB82], U-Net [GERLA83a], Token-Less Protocols (TLP) [RODR84], MAP [MARS82], and Buzznet [GERLS3b].

While the field of local area networks is usually meant to refer to packet switching as those described above, one is hard pressed not to ignore private branch exchange systems which are based on circuit switching and have originally been devised for voice communications, as these can be used for local computer communications.

1.2 Integrated-Services

When dealing with applications other than terminal access and computer-to-computer communications, there is often the need to satisfy stringent delay requirements. Applications which impose such requirements are referred to as "real-time" or "time-constrained" applications. There are many examples. Packetized voice in which voice is digitized, packetized at the source, transmitted over the network, and reconstructed at the destination is one such example. Excessive delays are disruptive to interactive speech conversations, and thus each bit of data must be delivered to the destination within a certain maximum delay. Another example is real-time process control in the factory, in avionics, etc. While voice and process control are both real-time applications, their attributes are different, and may often have to be handled differently. Consider for example uncompressed PCM speech. A voice source is synchronous in that it generates bits at a constant rate (64 Kbps), all bits are to be delivered within the same delay constraint, and a certain level of data loss is acceptable. In process control, sources typically require low throughput, are usually bursty, but delay constraints may be more severe, and loss may have disastrous effects.

So far, the study of integrated-services has been limited primarily to voice and data. In recent dissertation by Gonsalves [GONS86] an excellent overview of previous work on voice-data integration on local-area networks has been provided. What is most essential to note in all the schemes proposed for the integration of voice and data is the fact that they all take advantage of special characteristics of voice traffic. Typically vocoders digitize voice at a constant rate. Bits are grouped into packets which are then transmitted via the network to the destination vocoder. Accordingly packets are generated deterministically, i.e., at regular intervals.

Several papers addressed the performance (either simulated or measured) of

CSMA/CD networks. Among these we cite: [NUTTS2, TOBA82, GONS83, MUSS83, DETR84]. Others investigated the use of prioritized versions of CSMA-CD for the integration of voice and data, with voice assigned a higher priority than data [CHLA83, CHLA85, IIDA80, JOHN81, MAXE82, TOBA80, TOBA82]. Yet others have dealt with voice/data traffic on DAMA Networks, Fasnet and Expressnet in particular [LIMB82, LIMB83, FRAT81, TOBA83, FINE85a, FINE85b].

In sum, all previous studies have either evaluated the performance of various local area networks in view of the above stated requirements or devised new schemes which take advantage of the specific characteristic stated above, (i.e., deterministic packet generation) and are thus adequate for the integration of such applications.

Perhaps the only work so far which addresses the problem of real time applications in LANs is that by Kurose and Schwartz [KURO83] for CSMA networks. They have shown that for time-constrained communication over CSMA channels, the ordering imposed on message transmissions by the channel access strategy has a great effect on the distribution of message delay. In their work they have considered a generalized version of the time-based window algorithm (WCSMA) [GALL78, TOWS82] which allows to adaptively vary an imposed message transmission scheduling discipline in response to changing system demands. That is, while in WCSMA the objective is FCFS, here the distribution of packet delay is the critical performance measure, and the objective may for example be to maximize the percentage of messages with waiting times below a specified time bound.

1.1.3 Contributions

It is clear from the above survey of previous work that the integration of services on Local Area Networks had not been fully investigated. Indeed, most of the

studies concerned themselves with voice/data integration taking advantage of the deterministic nature of voice traffic, and thus are not applicable to the problem of integrated services in its general form. The work on time-constrained applications by Kurose and Schwartz is based on CSMA type networks and thus is not applicable to high speed operation of LANs, as CSMA performs particularly poorly in networks with high bandwidth.

The integration of services on Local Area Networks will undoubtedly require high speed operation. Indeed, with the multitude of applications that need to be integrated in LANs, bandwidth and delay requirements far exceed the capabilities of such LANs as a 10 Mb/s Ethernet. A single digitized video channel may require as high as 50 Mb/s. Graphic and image applications may too be high bandwidth applications. Accordingly, the solution to integrated services Local Area Networks will more naturally be found in LANs of the attempt-and-defer DAMA type such as Expressnet and Fasnet, rather than in the CSMA contention type. Moreover, to realize high speed operation of these LANs, the medium will have to be fiber optics, a unidirectional medium to which schemes such as Fasnet and Expressnet are appropriate.

The contributions of this thesis are in two areas. The first consists of new access methods along with transmission scheduling policies for the integration of services with differing requirements in unidirectional networks based on the attempt-and-defer access mechanism. The second is the study of fiber optics topological configurations of such unidirectional networks, in view of supporting large number of stations.

Work in the first area is accomplished in two steps. The first step, which constitutes Chapter 2, consists of a study of transmission scheduling policies that can be used in broadcast local area networks. More precisely, the problem is defined as

follows: each message is assumed to be determined by a delay-cost function, and the objective is to devise a policy which minimizes the expected cost per message. Six policies are considered in this chapter. The first four, first-come-first-served (FCFS), round robin (RR), highest-penalty-first-served (HPFS), and lowest-slack-first-served (LSFS), are well known. The other two are newly proposed in this dissertation and are referred to "earliest-slot-first-served" (ESFS) and "dynamic priority" (DP) scheduling policies. ESFS is applicable to fixed length single-packet messages and step delay-cost functions for messages. Due to the assumption of fixed-length messages, the channel time axis may be considered to be slotted. Under ESFS, messages are assigned to the slots in such a way that the total cost is minimized. Then the earliest scheduled message is selected for transmission. Dynamic-priority scheduling policy is applicable to general message-length distribution and general cost functions. Under this policy, to each message, a dynamic priority equal to its discounted cost rate divided by its critical remaining transmission time is assigned. When the channel becomes idle, the message with the highest priority is transmitted. These various schemes have been evaluated by simulation and the results have shown that the new proposed schemes outperform all previously known one.

Throughout Chapter 2, the distributed aspects of the scheduling problem are ignored, i.e., it is assumed that the propagation delay between every pair of stations is zero, and that every station has complete knowledge about all messages in the network. With these assumptions, the network could be viewed as a single-server queueing system. In chapter 3, the distributed implementation of the transmission scheduling policies is taken into account. More precisely, we are concerned with the incorporation of scheduling policies in broadcast local area networks in view of integrating services. There are two issues involved. One is the provision of efficient channel access schemes for broadcast local area networks operating at high speed. The other issue is the provision of necessary mechanisms for the implementation

of scheduling policies. As stated before, there are efficient channel access schemes based on attempt-and-defer access mechanisms which solve the first problem and which constitute the basis for our solution. With respect to the second problem, we choose to consider a distributed implementation for two reasons. The first is reliability: with distributed scheduling, the system is not dependent on the proper operation of a central scheduler. The second is improved performance, in terms of both delay and channel utilization. This is due to the fact that, in distributed scheduling, only stations' local information must be broadcast on the network to the central scheduler, and the control information from the central scheduler back to the stations is not required. We propose in this chapter a scheme which included (i) a mechanism for the dissemination of all relevant information required by the scheduling policy, and (ii) a station architecture which allows the implementations of any of the scheduling policies described in the previous chapter.

In chapter 4, we study fiber optic configurations for local area networks. Indeed, the medium that provides high bandwidth in local area networks is naturally fiber optics. While this technology lends itself easily to networks with point-to-point links, e.g. rings, the problem is more complex with broadcast bus networks. Due to intrinsic characteristics of certain fiber-optic components, in particular the reciprocity of couplers and the low impedance of detectors, to use this technology in the implementation of a shared multitapped bus requires careful design consideration. In this chapter we first examine the unidirectional networks in question (folded bus and dual bus topologies) and extract features which prove useful in devising fiber optics configurations. We then present various fiber optics configurations for their implementation. We then provide a complete power budget analysis of fiber optic configurations. This leads to the performance evaluation of the proposed configurations, in terms of number of stations that can be supported for a given power margin. The numerical results have shown that for a power margin of 40 dB and a coupler

excess loss of 1 dB, the number of stations ranges from 30 with reciprocal couplers to 70 with nonreciprocal ones, using a newly proposed "bypass" configuration.

Chapter 2

Transmission Scheduling Policies for Broadcast Local Area Networks

In a broadcast LAN, each message is received by all stations, regardless of its destination. A station discards received messages that are not intended for it. In such networks, whether implemented as a bus or a ring, only one message may be transmitted at a time. If the network supports integrated services, the most urgent message should be transmitted first. The urgency of a message may depend both on the application and on the generation time of the message. In this chapter, we are concerned with policies for ordering message transmissions on a broadcast channel. We refer to these as *transmission scheduling policies*. Throughout the chapter, we ignore the distributed aspects of the problem; i.e., we assume that the propagation delay between every pair of stations is zero, and that every station has complete knowledge about all messages in the network. With these assumptions, the network may be viewed as a single-server queueing system. The distributed implementation of the transmission scheduling policies is the subject of the next chapter. In the next section, the transmission scheduling problem is discussed in more detail. In section 2.2, a number of basic scheduling policies are reviewed, and

some new ones are proposed. In section 2.3, based on results obtained through regenerative simulation, the relative performance of these policies is evaluated.

2.1 The Transmission Scheduling Problem

A broadcast channel with zero propagation delay is equivalent to a single server, since only one message can be transmitted at any given time. All pending messages must therefore be queued, and a scheduling policy is used to select one of them for transmission, whenever the broadcast channel is idle. We define the delay of a message to be the time elapsed from its generation until the completion of its transmission on the channel. Based on the performance specifications of each application, we define a cost function for messages belonging to that application. (As examples, the cost functions appropriate for voice and data traffic are discussed below.) Let D_n and $c_n(\cdot)$ denote the delay and cost function of the n th message generated in the system, respectively. The average cost per message is given by

$$\bar{c} = \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{k=1}^n c_k(D_k). \quad (2.1)$$

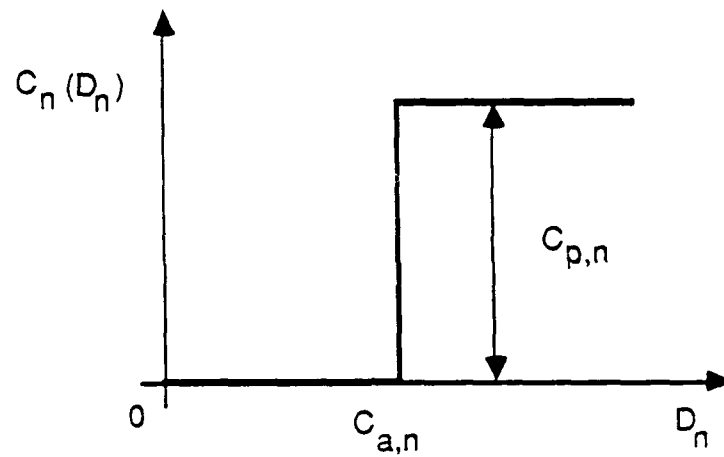
The objective is to find a scheduling policy which minimizes $\bar{c}$.

In scheduling theory, problems are classified according to the arrival process of jobs and their service-time distribution [FREN82]. When the number of jobs and their generation times are fixed and known a priori, the problem is *static*; when the jobs are generated randomly over a period of time, the problem is called *dynamic*. If the service times are known, the problem is called *deterministic*; if the service times are random, the problem is called *stochastic*. Based on this terminology, the scheduling problem formulated here is dynamic and stochastic. For a

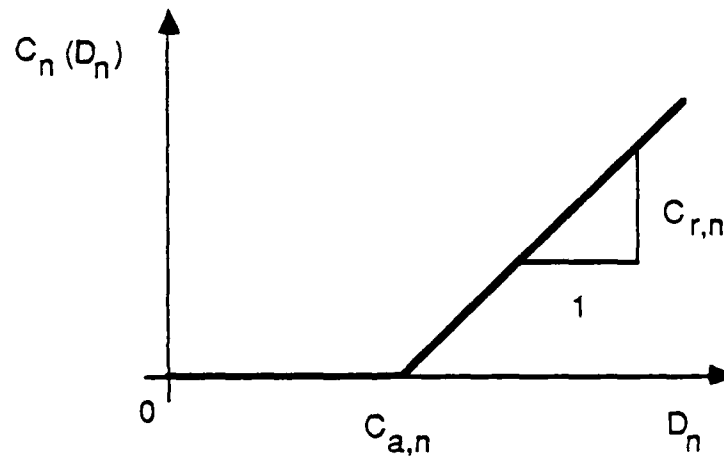
general service-time distribution and a general set of cost functions, a dynamic and stochastic scheduling problem, such as the one formulated here, is analytically intractable [RINN75]. (Below, we will also present a scheduling policy for a static and deterministic problem.) We will consider a number of suboptimal scheduling policies and compare their performance through simulation.

Before discussing the scheduling policies, let us examine a set of cost functions appropriate for voice and data traffic. In the case of interactive voice traffic, there exists a maximum allowable end-to-end delay. If the delay of a packet exceeds this maximum, it becomes out-dated and is lost. Depending on the desired quality and type of coding used, a certain amount of packet loss is tolerable. Based on this performance specification, the cost function appropriate for voice is a step function. As shown in Fig. 2.1(a), a step cost function, say that of the n th message generated in the system, is defined by two parameters: a *delay allowance* denoted by $C_{a,n}$, and a *penalty* denoted by $C_{p,n}$. The *deadline* d_n is defined as the time by which if the message is entirely transmitted, no penalty is incurred; deadline is equal to the arrival time plus the delay allowance. Under an optimal scheduling policy, increasing the penalty of voice messages effectively increases their priority over other messages and thus reduces the percentage of the messages lost.

For data traffic, the performance is specified in terms of some measure defined on the distribution of message delays. For example, the performance measure may be the average message delay. In that case a linear cost function should be used. The cost rate (i.e., the cost incurred per unit time) is chosen based on the priority of data traffic over other traffic. More generally, the performance measure may be the average delay (or some function thereof) of those messages whose delay exceeds a certain allowance time. The cost function used in that case would be a ramp function. As shown in Fig. 2.1(b), a ramp cost function, say that of the n th



(a) STEP COST FUNCTION



(b) RAMP COST FUNCTION

Fig. 2.1 Examples of delay cost function for the n th message generated in the system.

message generated in the system, is defined by two parameters: a delay allowance $C_{a,n}$, and a *penalty rate* denoted by $C_{r,n}$. As in the step cost function, the deadline is equal to the arrival time plus the delay allowance. However, unlike the step cost function, the messages which are not transmitted before their deadlines are not dropped.

2.2 Scheduling Policies

2.2.1 Basic Scheduling Policies

Four basic scheduling policies are considered here.

First-Come-First-Served (FCFS). Under this policy, the messages are transmitted in the order of their arrival. More specifically, a newly generated message is put at the end of the global queue. When the bus becomes idle, a packet belonging to the message at the head of the queue is transmitted. In the case of a multi-packet message, all the packets are transmitted back-to-back without interruption: i.e., there is no message preemption. The FCFS can be applied under any message-length distribution and any set of cost functions.

Round Robin (RR). As before, a newly generated message is put at the end of the (global) queue. When the bus becomes idle, a packet belonging to the message at the head of the queue is transmitted. In the case of a multi-packet message, after the transmission of each packet, the message is put at the end of the queue. This process is repeated until every packet in the message is transmitted. Like FCFS, the round-robin scheduling policy can be applied under any message-length distribution and any set of cost functions.

Highest-Penalty-First-Served (HPFS). As the name indicates, whenever the bus becomes idle, a packet belonging to the message with the highest penalty is transmitted. This implies that a penalty must be defined for every message. In the case of a message with a step cost function, the height of the step (i.e., the cost of missing the deadline) can be used as the penalty (see Fig. 2.1). Also, in the case of a message with ramp cost function, the slope of ramp (i.e., the cost rate after deadline) can be used as the penalty. However, in the case of a message with a general cost function, there is no natural choice for penalty. Hence, while the HPFS scheduling policy can be applied under any message-length distribution, it is somewhat restrictive on the set of cost functions.

Lowest-Slack-time-First-Served (LSFS). The *slack time* of a message at time t is defined as the remaining time to its deadline minus its remaining transmission time at t . Under the LSFS policy, whenever the bus becomes idle, a packet belonging to the message with the lowest slack time is transmitted. Clearly, when the slack time of a message becomes negative, it is no longer possible to transmit it before its deadline. If such a message has a step cost function, it must be dropped from the queue, since transmitting it would not reduce the incurred cost. If the message has some other cost function, say the ramp cost function, then it still must be transmitted. Therefore, at any point in time, some of the messages in the global queue may have negative slack times.

2.2.2 A Policy based on Static Scheduling

The scheduling problem, which was formulated in section 2.1, is both dynamic and stochastic, in that the messages are generated according to a random process and their lengths are randomly distributed. Since finding the optimal policy for such a scheduling problem is analytically intractable, any approach leading to a

suboptimal solution will be attractive. The approach which is taken here is as follows. Whenever a message must be selected for transmission, a schedule for transmitting all messages in the global queue is generated (the queue is rearranged). Then, the message with the earliest scheduled transmission time (head of queue), is selected. The objective in generating this schedule is to minimize the total cost incurred in delaying the messages which are already in the queue, ignoring any messages generated in future. Since, at any point in time, the number of messages in the global queue and their lengths and cost functions are known, this scheduling problem is a static and deterministic one.

For a general distribution of message lengths and a general set of cost functions, the static scheduling problem is NP complete in n , the number of queued messages [RINN75]; in fact, the optimal schedule can only be found through exhaustive search. For a very special case of this problem, we present an algorithm which is $O(n^2)$. Specifically, we assume that the messages are of fixed length, equal to that of a packet. Furthermore, we assume that all cost functions are step functions, though not necessarily with equal parameters values.

Due to the assumption of fixed-length single-packet messages, the channel time axis may be considered to be slotted, with the slot size equal to the transmission period of a message. The scheduling policy at the beginning of each slot selects a message for transmission. Initially, all messages are scheduled in the current and all future slots; and then, the message scheduled in the earliest slot is transmitted in the current slot. This scheduling policy is referred to as the *Earliest-Slot-First-Served* (ESFS) scheduling policy.

Now, we describe the algorithm used for scheduling messages in slots. Note that the objective of this (static) scheduling is to assign slots to queued messages to minimize the total penalty, incurred in dropping messages. Clearly, to minimize

the total penalty, a message should be dropped only when all the slots before its deadline are assigned to other messages with equal or higher penalty, and none of these slots can be freed by rescheduling the corresponding messages in some later slot. Based on this observation, a simple algorithm for optimal static scheduling is as follows.

Algorithm 2.1: Messages are first arranged in order of decreasing penalties. Each message is then scheduled in the latest unassigned slot before its deadline. If there are no unassigned slots before its deadline, the message is dropped.

Before showing the optimality of this algorithm, let us illustrate its operation through the example shown in Table 2.1. Without loss of generality, we are considering the schedule generated at the start of slot 0. Each message is denoted by a pair (p, d) , where p is the value of its penalty and d is the latest slot before its deadline. $\mathcal{B}$ is the set of backlogged messages at the start of slot 0. Each iteration corresponds to scheduling a single message. In iterations 1 through 3, messages $(5,6)$, $(5,2)$, and $(4,1)$ are scheduled in slots 6, 2 and 1, respectively. In iteration 4 of the algorithm, since slots 1 and 2 are already assigned, message $(3,2)$ is scheduled in slot 0. In iteration 5, the message $(2,6)$ is scheduled in slot 5. In iteration 6, since slot 0 is already assigned, message $(2,0)$ is dropped. In the last iteration $(2,3)$ is assigned in slot 3. Based on the resulting schedule $(3,2)$ is transmitted in slot 0.

To prove that Algorithm 2.1 is optimal, we show that if a message is dropped, it could not have been scheduled without dropping another message with a greater or equal penalty. Let (p, d) be the message that is dropped at some iteration. Let s be the earliest unassigned slot at this iteration. Clearly, s is greater than d , otherwise (p, d) could have been scheduled in slot s . Moreover, since according to the algorithm every message is scheduled in the latest unassigned slot before its deadline, all messages which are already scheduled in the slots before slot s must

$$\mathcal{B} = \{(3,2), (2,6), (5,6), (2,0), (5,2), (2,3), (4,1)\}$$

Iteration	Slots						
	0	1	2	3	4	5	6
1							(5,6)
2			(5,2)				(5,6)
3		(4,1)	(5,2)				(5,6)
4	(3,2)	(4,1)	(5,2)				(5,6)
5	(3,2)	(4,1)	(5,2)			(2,6)	(5,6)
6	(3,2)	(4,1)	(5,2)			(2,6)	(5,6)
7	(3,2)	(4,1)	(5,2)	(2,3)		(2,6)	(5,6)

Table 2.1 Example illustrating the operation of algorithm 2.1. Each message is denoted by a pair (p, d) , where p is the value of its penalty and d is the latest slot before its deadline. $\mathcal{B}$ is the set of backlogged messages at the start of slot 0. Every iteration corresponds to scheduling a single message.

have a deadline less than s . Also, since the scheduling is performed in the order of decreasing penalties, all these messages must have a penalty greater than or equal to p .

We compute the computational complexity of the algorithm as follows. In each iteration of Algorithm 2.1, a single message is scheduled. Before the i th iteration of the algorithm, there are at most $i - 1$ assigned slots. Therefore, finding the latest unassigned slot prior to the deadline of the message being scheduled, involves

$O(i)$ steps. If there are n backlogged messages, the total number of steps is thus $O(n^2)$. (Note that sorting by penalty is $O(n \log n)$ and hence does not increase the complexity.)

We next present a slight modification to algorithm 2.1, which obviates the need for sorting the messages prior to scheduling. This will prove useful in more realistic situation where new messages are being generated all time. The modified algorithm is as follows.

Algorithm 2.2: Each message is scheduled in the latest slot before its deadline which has not already been assigned to a message with equal or higher penalty. If there is no such slot, the message is dropped. If the message is scheduled in a slot which has previously been assigned to a message with a lower penalty, then that message is removed and rescheduled, as just described. The algorithm terminates when every message is either scheduled or dropped.

Observe that a removed message can only be rescheduled in the slots before the slot from which it was removed. Also, since a message cannot cause the removal of a higher penalty message, the above algorithm eventually terminates. Since every removed message is either dropped or rescheduled to an earlier slot than the one from which it was removed, the algorithm eventually terminates. The operation of Algorithm 2.2 for our example is illustrated in Table 2.2. In iterations 1 and 2, messages (3,2) and (2,6) are scheduled in their deadline slots, namely slots 2 and 6, respectively. In iteration 3, message (5,6) bumps message (2,6) from slot 6 to slot 5. In iteration 4, message (2,0) is scheduled in slot 0. In iteration 5, message (5,2) bumps message (3,2) from slot 2 to slot 1. In iteration 6, message (2,3) is scheduled in its deadline slot. In the final iteration, message (4,1) bumps message (3,2) from slot 1 to slot 2, and message (3,2), in turn, bumps message (2,0). Message (2,0) is then dropped. As in Algorithm 2.1, message (3,2) is transmitted in slot 0.

$$\mathcal{B} = \{(3,2), (2,6), (5,6), (2,0), (5,2), (2,3), (4,1)\}$$

Iteration	Slots						
	0	1	2	3	4	5	6
1			(3,2)				
2			(3,2)				(2,6)
3			(3,2)			(2,6)	(5,6)
4	(2,0)		(3,2)			(2,6)	(5,6)
5	(2,0)	(3,2)	(5,2)			(2,6)	(5,6)
6	(2,0)	(3,2)	(5,2)	(2,3)		(2,6)	(5,6)
7	(3,2)	(4,1)	(5,2)	(2,3)		(2,6)	(5,6)

Table 2.2 Example illustrating the operation of Algorithm 2.2. Every message is denoted by a pair of integers (p, d) , where p is its penalty and d is the latest slot before its deadline. $\mathcal{B}$ is the set of backlogged messages at the start of slot 0. Every iteration corresponds to scheduling a single message.

Algorithms 2.1 and 2.2 described above are optimum for a static scheduling problem. When used for a dynamic scheduling problem, they are no longer optimum. For example, there may be new arrivals during slot 0 which, when taken into account, may result in the preemption and loss of message $(4,1)$. The expected loss would be less if message $(4,1)$ is transmitted in slot 0, and message $(3,2)$ is scheduled in slot 1. In general, an improvement in performance is possible if one permutes the schedule in such a way that the messages with higher penalty are transmitted first, while none of the lower penalty messages are dropped. (For the traffic

model considered in the next section. this permutation did not result a noticeable improvement, and thus is not included in the performance results.)

An advantage of Algorithm 2.2 over Algorithm 2.1 is that the schedule generated for one slot can be incrementally updated for the next slot. The incremental changes in the schedule consist of removing the transmitted message from the queue and scheduling newly generated messages.

2.2.3 Dynamic-Priority Scheduling Policy

Here, we propose a new priority rule for the transmission scheduling problem. which can be viewed as a generalization of the $C\mu$ rule. (C refers to cost per unit time, and μ refers to service rate.) The $C\mu$ rule which is due to Klimov [KLIM74, KLIM 78], is the optimal solution to a dynamic and stochastic scheduling problem. under a number of restrictive assumptions on the message-length distributions and cost functions. The new policy, on the other hand, while not optimal, can be applied to general message-length distributions and general cost functions. Also, in the $C\mu$ rule, the message priorities are static, while in the new policy, the priorities are dynamic. We refer to the proposed scheduling policy as the *Dynamic-Priority (DP) Scheduling Policy*.

Before presenting the new scheduling policy, we review the $C\mu$ rule in the context of the transmission scheduling problem. The messages are assumed to be generated according to a Poisson process. Every message belongs to one of K classes, which are denoted by $1, 2, \dots, K$. Messages belonging to the same class have the same cost function and the same length distribution. The cost functions are assumed to be linear and we use C_k denote the cost rate of messages in class k , for $k = 1, 2, \dots, K$. It is assumed that the transmission times of messages in class k is exponentially distributed with mean $\bar{T}_k$. Under the $C\mu$ rule, a static priority equal

to $C_k \mu_k$ is assigned to every message in class k ; the priorities are used to transmit the messages according to a preemptive discipline: i.e., during the transmission of a message, if another message with a higher priority is generated, the transmission of the former will be interrupted and that of the latter will be initiated. Note that, since the message length distributions are exponential and thus possess the memoryless property, it makes no difference whether the transmission of the preempted message is resumed from where it was interrupted, or from the beginning. Also note that, although the lengths of messages are known after their generation, this information is not used in $C\mu$ rule.

The fact that a static priority, as opposed to a dynamic one, is optimal for a dynamic and stochastic scheduling problem, is due to the following properties of the model. Firstly, the linear cost function satisfies the memoryless property: i.e., at any point in time, the marginal cost of delaying the message another unit of time does not depend on how long the message has already been delayed. Secondly, the exponential transmission-time distribution also satisfies the memoryless property: i.e., at any point in time, the distribution of remaining transmission time of a message, is independent of how long it has been in transmission.

Based on the above observations regarding the properties of the model, the form of static priority is intuitively appealing. Assume that, at time t , one of two messages, say M_1 and M_2 , is to be selected for transmission. The objective is to minimize the expected sum of their delay costs. Due to the memoryless properties of the linear cost function and the exponentially-distributed transmission time, we can assume that these messages were both generated in the system at time t . Let k_1 and k_2 denote the class of M_1 and M_2 , respectively. Then, if M_1 is transmitted at time t , the expected sum of costs will be $\bar{T}_{k_1} C_{k_1} + (\bar{T}_{k_1} + \bar{T}_{k_2}) C_{k_2}$. Otherwise, it will be $(\bar{T}_{k_2} + \bar{T}_{k_1}) C_{k_1} + \bar{T}_{k_2} C_{k_2}$. Message M_1 should be transmitted first if the

latter sum exceeds the former one, which is the case if $C_{k_1}/\bar{T}_{k_1} = C_{k_1}\mu_{k_1}$ is greater than $C_{k_2}/\bar{T}_{k_2} = C_{k_2}\mu_{k_2}$. In other words, the message with the highest static priority should be transmitted first.

Now, we derive the dynamic-priority scheduling policy, by relaxing the assumptions underlying the $C\mu$ rule. In fact, we make no assumptions regarding the message-arrival process, the distribution of message lengths, and the cost functions. But we allow message preemption only at packet boundaries.

Under the $C\mu$ rule, where the cost functions are assumed to be linear, the cost rate, C_k , provides a measure of urgency for messages in class k . To define the dynamic-priority policy, we need a similar measure of urgency for messages with general cost functions. Let M_n be the n th message generated in the system. Let τ_n and $c_n(\cdot)$ denote the generation time and cost function of M_n , respectively. The cost incurred by M_n if its transmission is completed at time t is $C_n(t - \tau_n)$. The instantaneous cost rate at time t is $c'(t - \tau_n)$, where $c'_n(t) = \frac{d}{dt}c_n(t)$. Clearly, $c'(t - \tau_n)$ cannot be used as the role measure of urgency for M_n at time t , since it does not account for possible abrupt cost jumps in the immediate future. Instead, we use a *discounted average* of the cost rate over all future time. Let $T_n(t)$ denote the transmission time of packets remaining in M_n at time t . We define the *urgency* of message M_n at time t as

$$C_n(t) \triangleq \int_{t+T_n(t)}^{\infty} \frac{e^{-x/\alpha}}{\alpha} c'(x - \tau_n) dx \quad (2.2)$$

where α is the *discounting factor*. Since the minimum possible delay of M_n is $t + T_n(t) - \tau_n$, the cost rates corresponding to delays less than this value are not included in the integral. We define the dynamic priority of message M_n as

$$P_n(t) \triangleq \frac{C_n(t)}{T_n(t)} \quad (2.3)$$

Under the dynamic-priority scheduling policy, when the channel is idle, the message with highest dynamic priority is transmitted. (The computational complexity is therefore linear in the number of backlogged messages.) Note that the static priority of $C\mu$ rule is equal to the (*constant*) cost rate divided by the *expected* remaining transmission time, whereas, the dynamic priority is equal to the *discounted* cost rate divided by the *actual* remaining transmission time.

The choice of the discounting factor, α , has a direct effect on the performance of the dynamic-priority policy. For very small values of α , only the cost incurred in the immediate future is effectively taken into account. This is not optimal, for example, in the case of step cost functions, it can cause the transmission of a message at the expense of dropping one with a later deadline but a higher penalty. For very large values of α , on the other hand, the cost too far into the future is taken to account and the policy becomes too conservative. In the case of step cost functions, a message may be transmitted at the cost of dropping one with an earlier deadline and lower penalty, where in fact both of them could have been transmitted if the latter were transmitted before the former. As we shall observe, under any traffic condition there is an optimal value of α at which the expected delay cost is minimized. We refer to the dynamic-priority policy with this optimal α as DPO policy. Since in an actual system it is not practical to optimize the discounting factor under different traffic conditions, we shall also consider the performance of the dynamic-priority policy with a fixed α . Here, we set α equal to the average transmission time of a message, and refer to the policy as DP1.

For step and ramp cost functions, we have derived closed-form formulas for the dynamic priority of a message. If message M_n has a step cost function, with penalty

of $C_{p,n}$ and deadline of d_n , its priority is given by

$$P_n(t) = \begin{cases} \frac{C_{p,n}}{\alpha T_n(t)} \exp[-(d_n - t - T_n(t))/\alpha], & \text{if } d_n - t - T_n(t) > 0; \\ 0, & \text{otherwise.} \end{cases} \quad (2.4)$$

If message M_n has a ramp cost function, with cost rate of $C_{r,n}$ and deadline of d_n , its priority is given by

$$P_n(t) = \begin{cases} \frac{C_{r,n}}{T_n(t)} \exp[-(d_n - t - T_n(t))/\alpha], & \text{if } d_n - t - T_n(t) > 0; \\ \frac{C_{r,n}}{T_n(t)}, & \text{otherwise.} \end{cases} \quad (2.5)$$

In order to gain some insight into the behavior of the dynamic-priority policy, we have plotted in Fig. 2.2 and Fig. 2.3, the urgency of message M_n as a function of the earliest possible completion of transmission (i.e., $t + T_n(t)$) for the two cases of step cost function and ramp cost function, respectively. The slack time at time t is $d_n - (t + T_n(t))$. In the case of a step cost, the urgency increases exponentially as the slack time approaches zero. When the slack time becomes negative, the urgency (and thus the priority) drops to zero. This is to be expected, since if slack time is negative, the message cannot possibly be transmitted before its deadline and might as well be dropped. In the case of ramp cost, again, the urgency increases exponentially as the slack time approaches zero. When the slack time becomes negative, the urgency becomes constant and equal to the cost rate $C_{r,n}$. The priority becomes $C_{r,n}/T_n(t)$, which is similar to the priority in $C\mu$ rule, namely $C_k/\bar{T}_k$.

2.3 Performance

The performance of the scheduling policies which were presented in the previous sections are evaluated here. The performance measure is the expected cost per

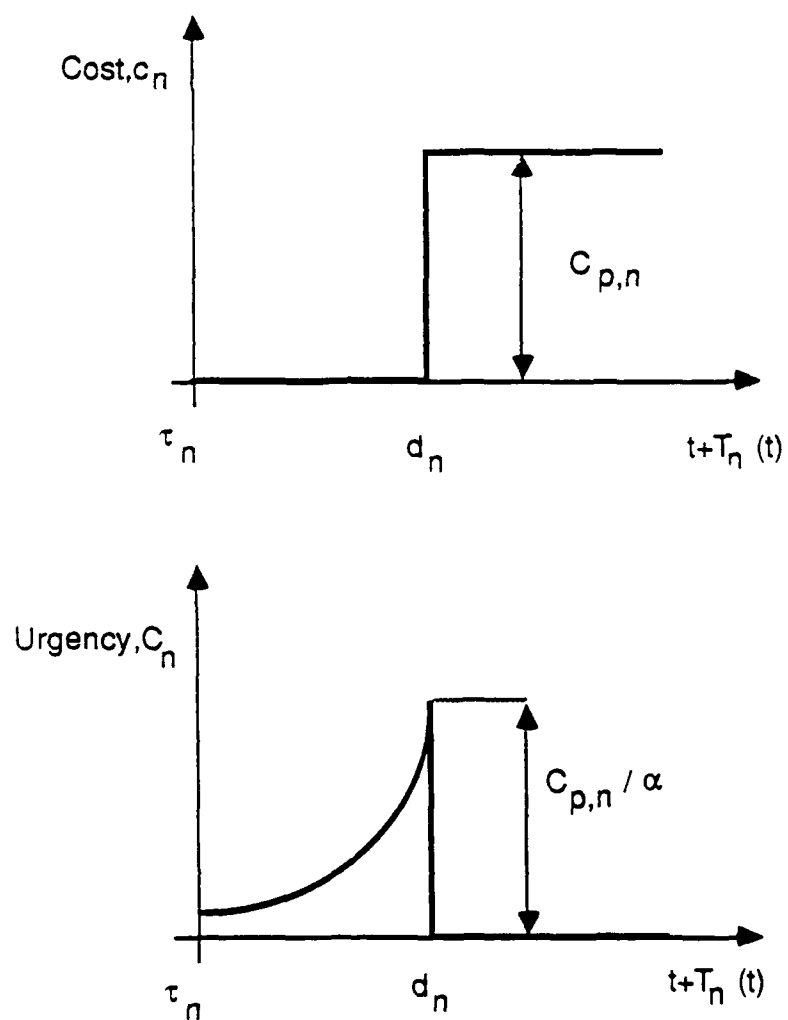


Fig. 2.2 Delay cost and urgency versus the earliest possible completion time for the step cost function.

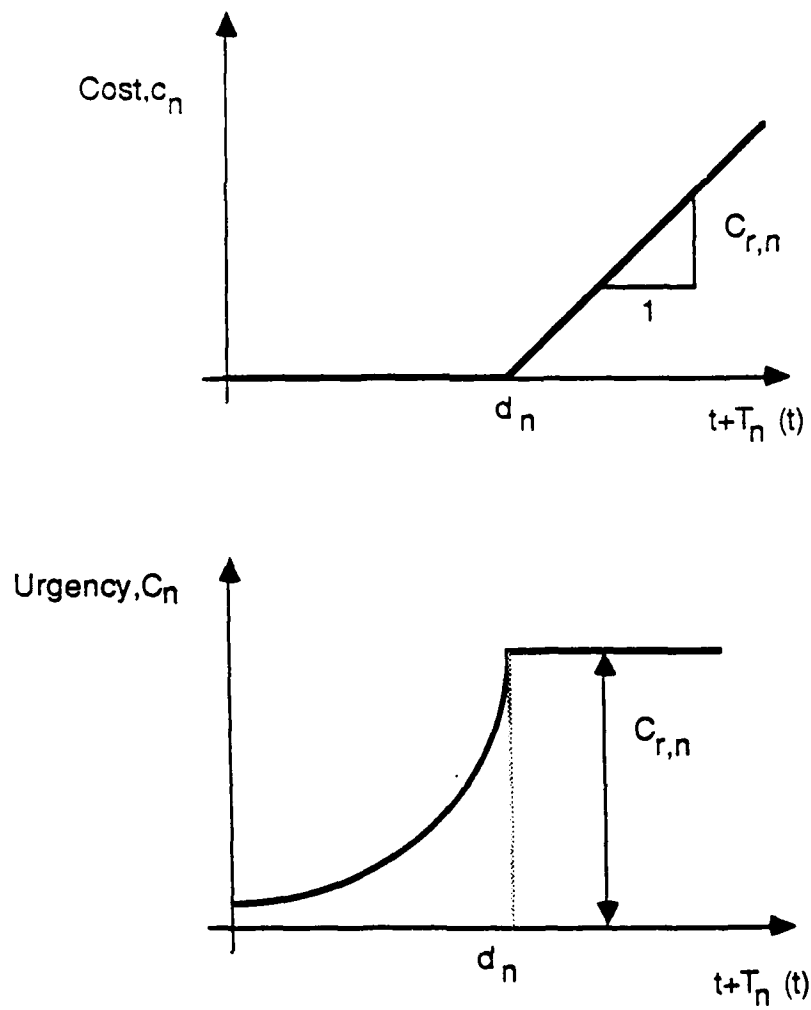


Fig. 2.3 Delay cost and urgency versus the earliest possible completion time for the ramp cost function.

message. As mathematical analysis proved to be complex, regenerative simulation is used to estimate the performance.

2.3.1 Traffic Model

In order to put the various policies to the test and to compare their performance, we choose the following traffic model. Messages are assumed to be generated according to a Poisson process. Two message length distributions are considered: deterministic and exponential. The length of message M_n is denoted by L_n , and its transmission time by T_n . The mean message length is denoted by $\bar{L}$, and the mean transmission time by $\bar{T}$. The maximum allowable packet length is denoted by K , and its transmission time by S . If L_n is not greater than K , message M_n is transmitted as a single packet. Otherwise, it is divided into a number of packets which, except for the last one, are all of length K . The channel load factor ρ is defined by $\rho: \rho \triangleq \lambda \bar{L} / R = \lambda \bar{T}$.

Two types of cost functions are considered: step and ramp. As in Fig. 2.1, the delay allowance of message M_n is denoted by $C_{a,n}$. Clearly, the delay allowance of a message must be greater than its transmission time. We choose an exponential distribution for the difference between the delay allowance and message transmission time. Let $\bar{C}_a$ denote the mean delay allowance. We also choose the penalties $C_{p,n}$ (for the step function), and $C_{r,n}$ (for the ramp function), to be exponentially distributed, with mean $\bar{C}_p$ and $\bar{C}_r$, respectively.

This traffic model is an arbitrary but general one. It introduces no bias in favor of any specific policy, and therefore is considered here to be a valid one for testing the various schemes.

2.3.2 Simulation

Discrete event simulation is used to evaluate the performance of scheduling policies. Based on the standard regeneration method [CRAN75], point and interval estimates of the expected cost per message are obtained.

The system state consists of the set of messages queued at the network, their arrival times, remaining transmission times, cost functions, and the identity of message being transmitted. The state in which there are no messages queued in the network and none is being transmitted, is referred to as the *empty-and-idle* state. Since messages are generated according to a Poisson process, the successive transitions to the empty-and-idle state constitute as set of regeneration points. The simulation is initiated from the empty-and-idle state.

The delay and cost function of the n th message generated in the system, M_n , are denoted by D_n and $c_n(\cdot)$, respectively. (If M_n is dropped, D_n will be equal to infinity.) The objective is to obtain point and interval estimates for the expected cost per message $\bar{c}$ as defined by (2.1). Let γ_k denote the number of messages generated in the k th regenerative cycle. Let $\beta_k \triangleq \gamma_1 + \gamma_2 + \dots + \gamma_k$ for all $k \geq 0$. Then, the total cost incurred for the messages generated in the k th cycle is given by

$$Y_k \triangleq \sum_{n=\beta_{k-1}+1}^{\beta_k} c_n(D_n) \quad (2.6)$$

Since the regeneration cycles are i.i.d, the pairs (Y_k, γ_k) for $k \geq 0$, are i.i.d.

For a given number of regeneration cycles say m , the point estimate of $\bar{c}$ is denoted by $\hat{c}(m)$ and is given by

$$\hat{c}(m) = \frac{\sum_{k=1}^m Y_k}{\sum_{k=1}^m \gamma_k} \quad (2.7)$$

Based on the law of large numbers for i.i.d. sequences of random variables, $\hat{c}(m)$ is a strongly consistent estimate of $\bar{c}$: i.e., it converges with probability one to $\bar{c}$ as $m \rightarrow \infty$.

Using the central limit theorem for i.i.d random variables, asymptotic confidence interval of $\bar{c}$ can be obtained [CRAN75]. Let $\hat{I}(m)$ be the 95% asymptotic confidence interval of $\bar{c}$, after m regeneration cycles. This means that, for large m , the interval $\hat{I}(m)$ contains the unknown constant $\bar{c}$ approximately 95% of the time. For the simulation results presented below, the number of regeneration cycles m is not fixed at the beginning of the simulation. For every estimate, enough regeneration cycles are obtained until half of the length of the 95% confidence interval $\hat{I}(m)$ becomes less than 10% of the value of point estimate $\hat{c}(m)$.

2.3.3 Numerical Results

The simulation results presented here provide a comparative performance evaluation of the scheduling policies which were presented in section 2. As stated above, the performance measure is the expected cost per message, $\bar{c}$. The scheduling policies along with their abbreviations are listed below.

FCFS	First Come First Served
RR	Round Robin
HPFS	Highest Penalty First Served
LSFS	Lowest Slack time First Served
ESFS	Earliest Slot First Served
DP1	Dynamic Priority with $\alpha = \bar{T}$
DPO	Dynamic Priority with Optimized α

Before comparing the scheduling policies, let us examine the effect of the discounting factor on the performance of the dynamic-priority policy. Figure 2.4 depicts $\bar{c}$ versus the normalized discounting factor, $\alpha/\bar{T}$, for different combinations of

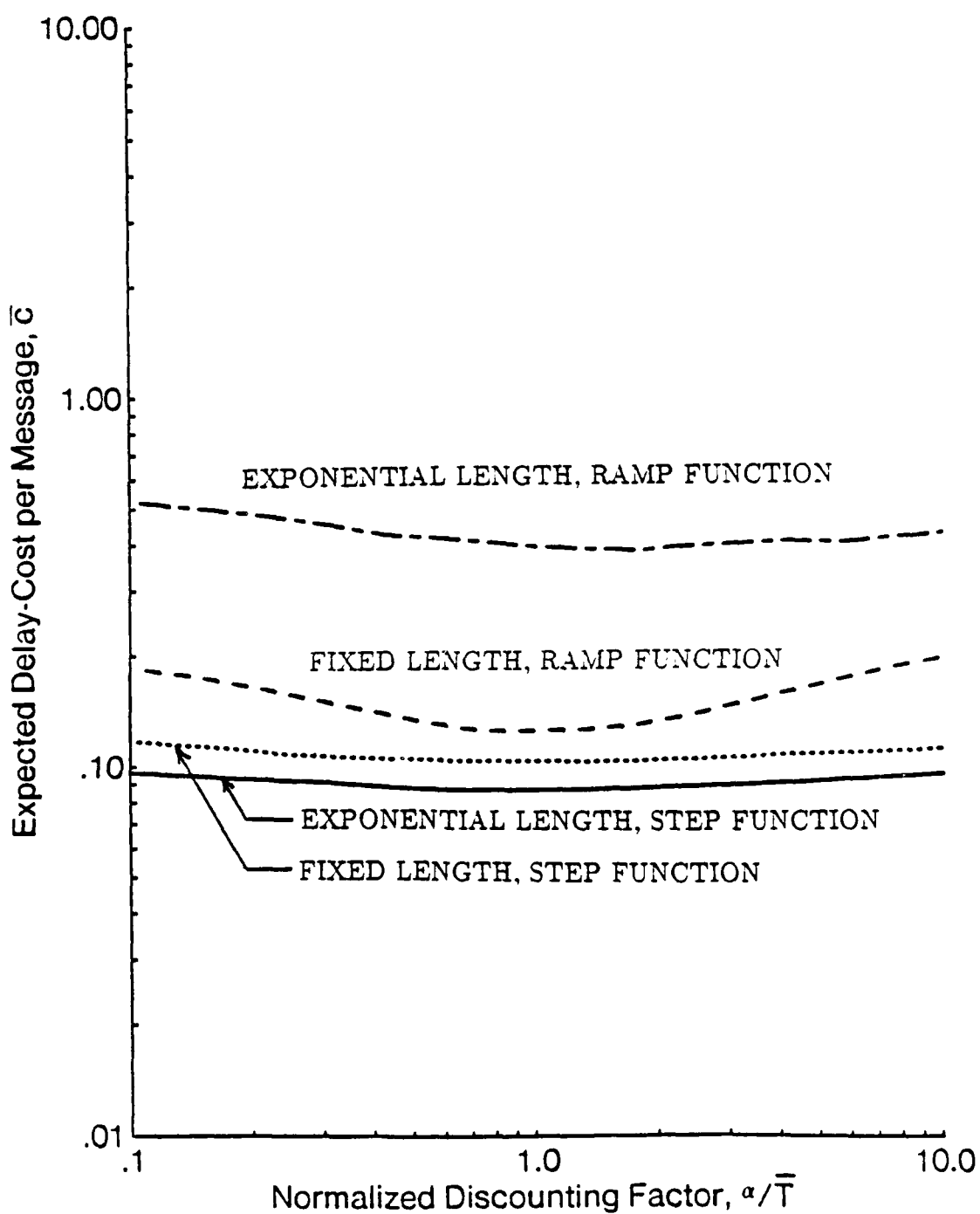


Fig. 2.4 Expected cost per message versus normalized discounting factor, for different combinations of message-length distributions and cost functions. $\rho = 0.8$, $\bar{L} = K$, $\bar{C}_a = 5\bar{T}$, and $\bar{C}_p = 1$.

message-length distributions and cost functions, $\rho = 0.8$, $\bar{L} = K$, $\bar{C}_a = 5\bar{T}$, and $\bar{C}_p = 1$. (Note that multiplying $\bar{C}_p$ by a factor will only scale $\bar{c}$ by that factor.) As shown in Fig. 2.4, for every combination of message-length distribution and cost function, there is an optimal α , at which the expected cost per message is minimum. Furthermore, $\bar{c}$ is not very sensitive to α around this minimum.

Figures 2.5-2.8 show the expected cost per message versus normalized mean delay allowance, $\bar{C}_a/\bar{T}$, for various combinations of message length distributions and cost functions. In these figures, $\rho = 0.8$, $\bar{L} = K$, and $\bar{C}_p = 1$. Clearly as the mean delay allowance increases, the expected cost decreases. In all cases the dynamic-priority policies perform better than all other policies. Furthermore, the ranking of the policies in terms of their performance is the same throughout. We note, however, that the difference in performance is more pronounced when the cost functions are of the ramp type (figures 2.6 and 2.8) than when they are of the step type (figures 2.5 and 2.7). This is explained as follows. With a step cost function, if a message is not transmitted before its deadline, the message is lost; although a penalty is incurred, the system is temporarily relieved, in that there is one less message to transmit. However, when a message has a ramp cost function, it must eventually be transmitted, even if it is past its deadline. This means that in the latter case, the load is consistently higher, causing inferior schemes to become even worse, causing in turn large gaps among the policies.

In figures 2.5 and 2.6, where the message length is deterministic (and contains a single packet, since $\bar{L}$ was chosen equal to K) FCFS and round robin are identical. This is not the case for random message lengths (Figures 2.7 and 2.8) although the difference between the two is minimal. FCFS and round robin perform the worst among all policies as they do not take into account information regarding deadlines and penalties; all other policies (LCFS, ESFS, and DP) perform better owing to

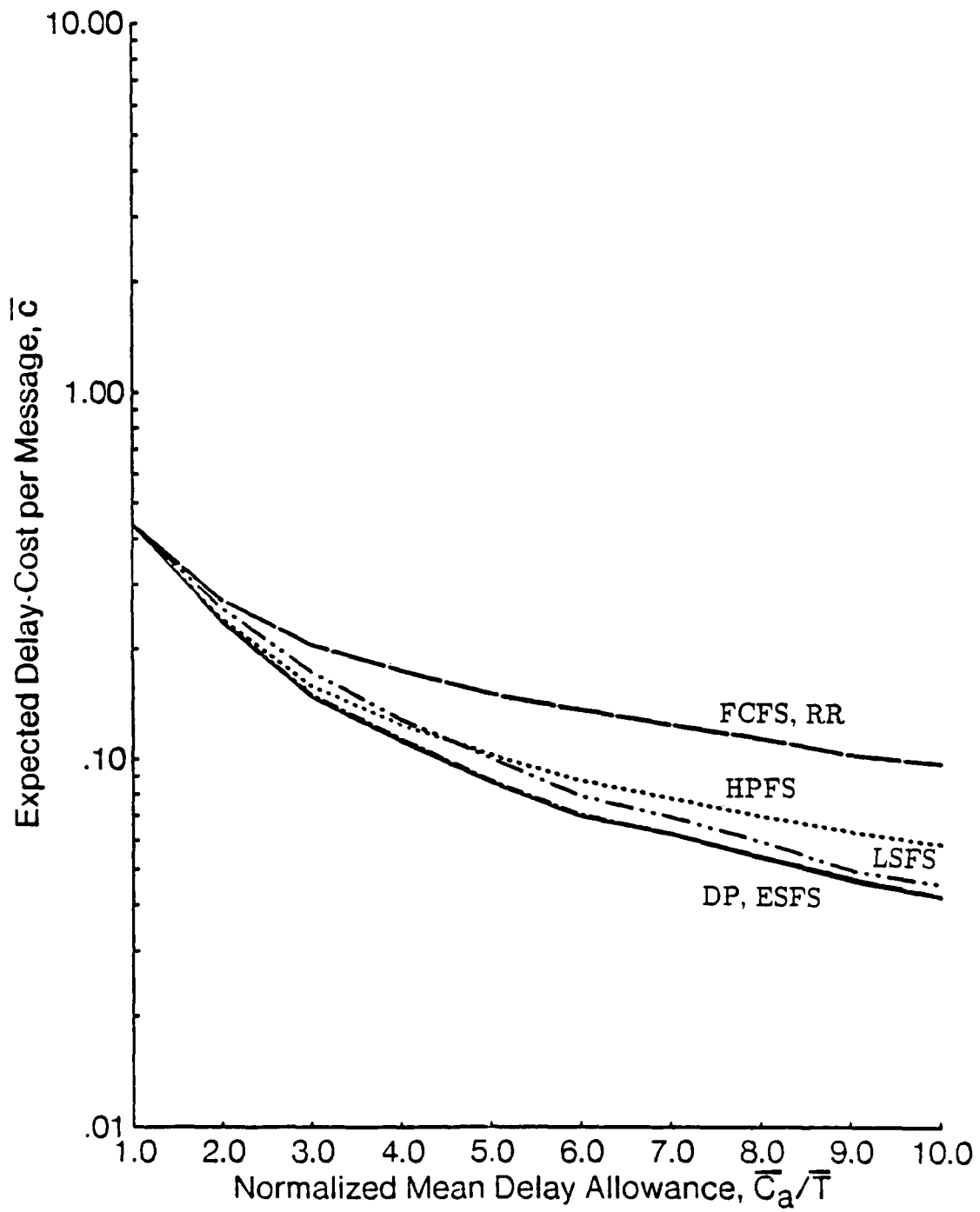


Fig. 2.5 Expected cost per message versus normalized mean delay allowance, for fixed-length messages, step cost functions, $\rho = 0.8$, $\bar{L} = K$, and $\bar{C}_p = 1$.

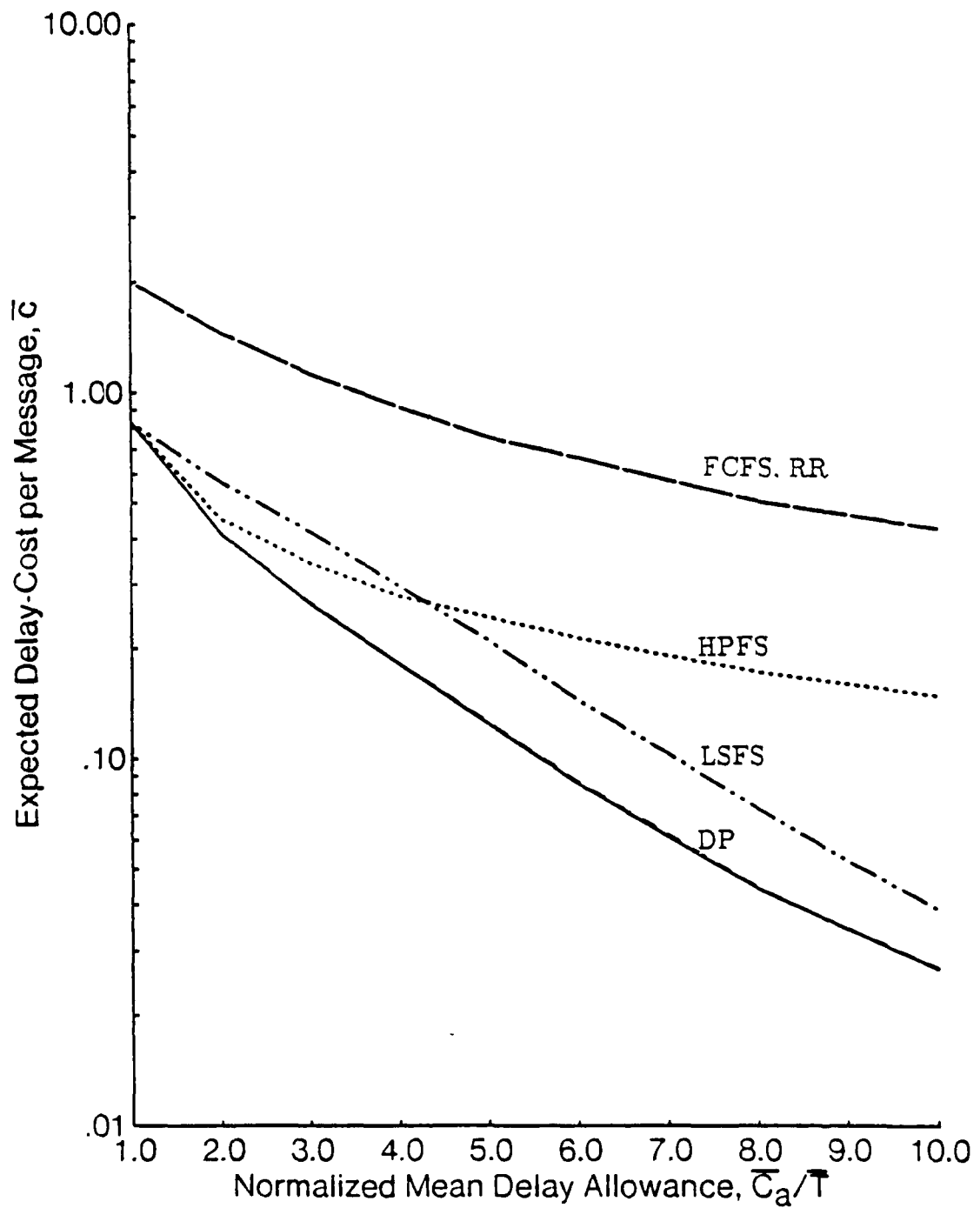


Fig. 2.6 Expected cost per message versus normalized mean delay allowance, for fixed-length messages, ramp cost functions, $\rho = 0.5$, $\bar{L} = K$, and $\bar{C}_p = 1$.

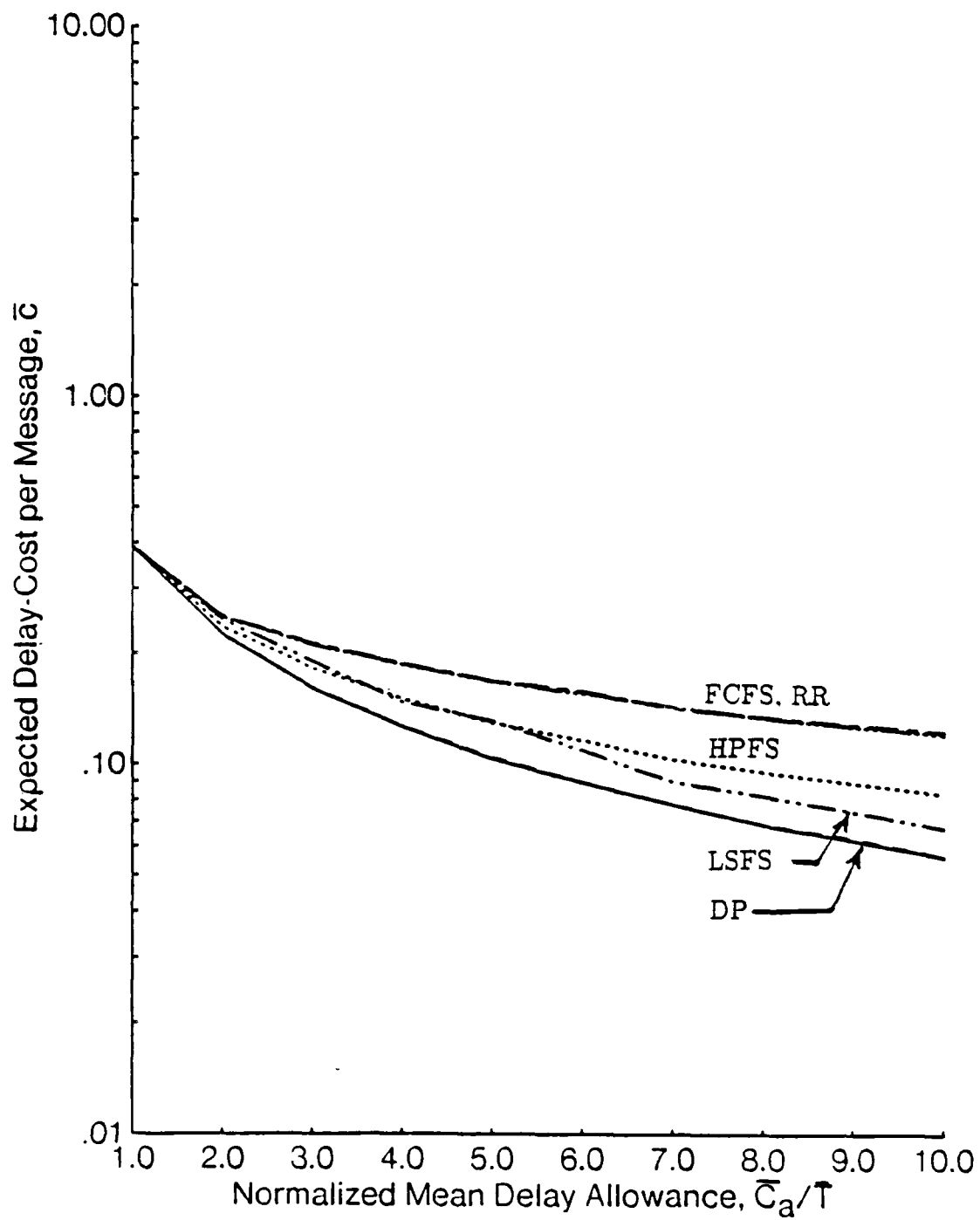


Fig. 2.7 Expected cost per message versus normalized mean delay allowance, for exponentially-distributed message lengths, step cost functions, $\rho = 0.5$, $\bar{L} = K$, and $\bar{C}_p = 1$.

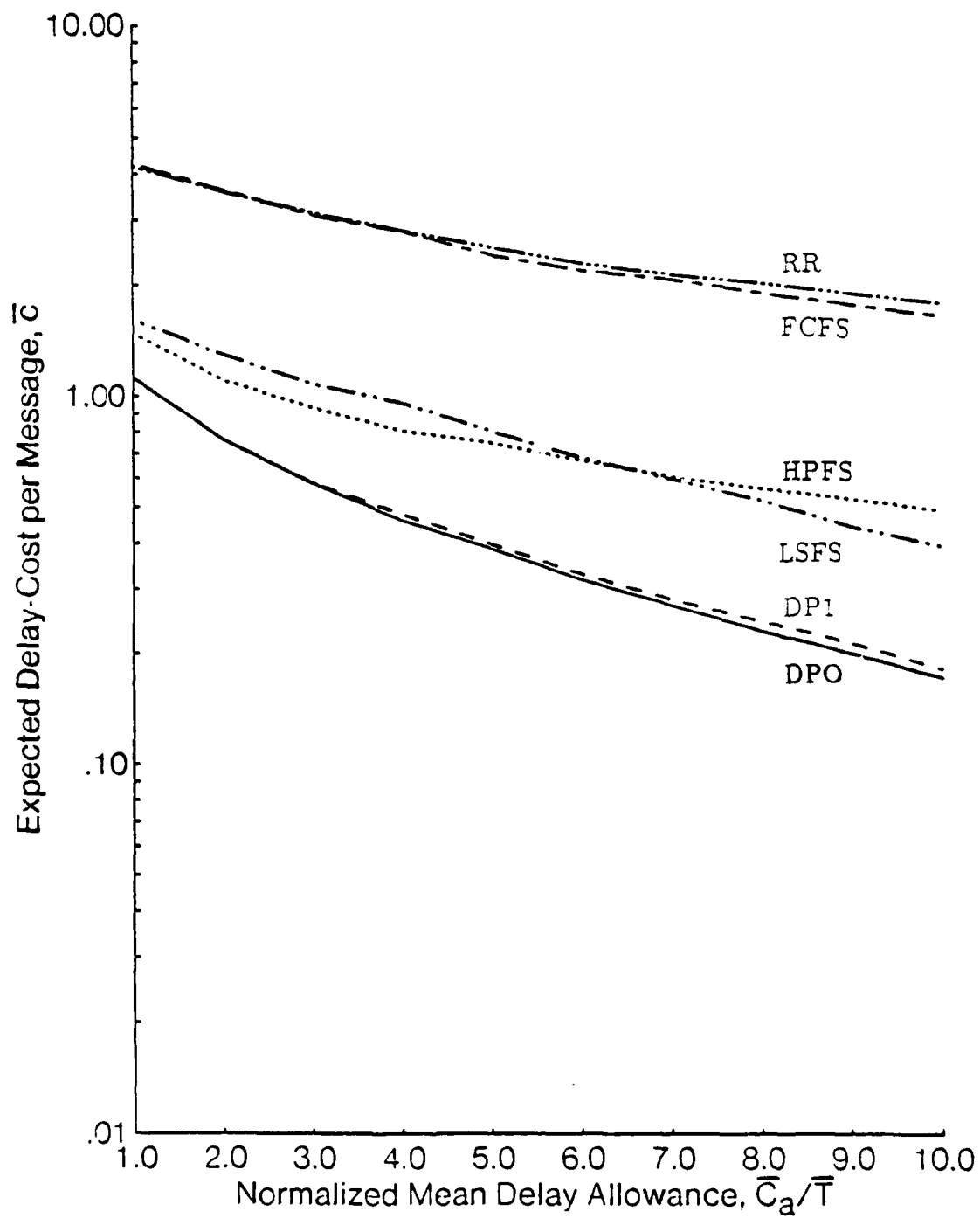


Fig. 2.8 Expected cost per message versus normalized mean delay allowance, for exponentially-distributed messages, ramp cost functions, $\rho = 0.3$, $\bar{L} = K$, and $\bar{C}_p = 1$.

using partial or complete information. Finally, it is interesting to note that, in the case of fixed message length (with $\bar{L} = K$) and step cost function (Fig. 2.5), when $\bar{C}_a/\bar{T} = 1$ only those messages which are generated in an empty system are transmitted and all other messages are lost, since in this case the delay allowance of every message is equal to its transmission time.

Figures 2.9-2.12 show the performance $\bar{c}$ versus ρ for the various combinations of message-length distributions and cost functions. In these figures, $\bar{C}_a/\bar{T} = 3$. Both DP1 and DPO outperform all other policies under all load conditions. The ranking among the policies is the same as in figures 2.5-2.8 (where $\rho = 0.5$) except for the crossover between HPFS and LSFS as ρ increases toward 1. The crossover is explained by the fact that, when ρ is large, the probability that a message exceeds its deadline becomes high, rendering the information regarding the slack time (used by LSFS) less important than the information regarding penalty and penalty rate (used by HPFS).

Figures 2.13-2.16 show the performance versus the normalized maximum packet length, $K/\bar{L}$. For fixed message lengths (figures 2.13 and 2.14), the expected cost $\bar{c}$ is constant for all values of $K/\bar{L} \geq 1$, since in this case the message is transmitted as a single packet (albeit partially full); the interesting part of these curves is therefore the range $K/\bar{L} < 1$. When $K/\bar{L} < 1$, a message contains more than one packet and thus RR and FCFS are not identical. In general, two effects come into play when messages contain more than one packet: on one hand, there are as many scheduling-decision opportunities as there are packets, and thus, the more packets there are, the more opportunities there are; on the other hand, service preemption tends to increase the mean message delay [KLEI76], and thus could drive the cost $\bar{c}$ up. For RR, which does not use any information regarding penalty and cost, only the second effect comes into play causing $\bar{c}$ to exceed that of FCFS and to increase

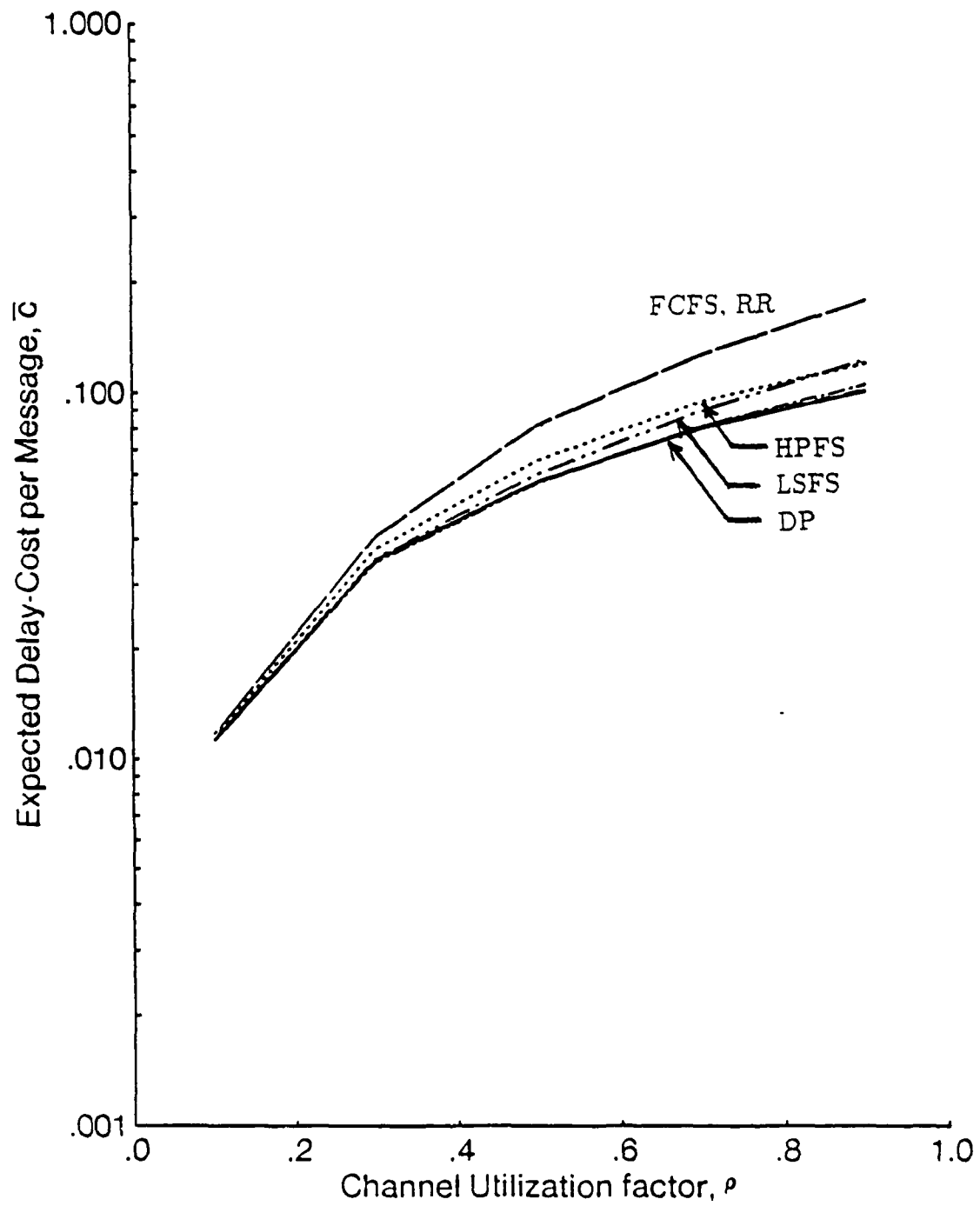


Fig. 2.9 Expected cost per message versus load factor ρ for fixed-length messages, step cost functions, $\bar{C}_a = 5\bar{T}$, $\bar{L} = K$, and $\bar{C}_p = 1$.

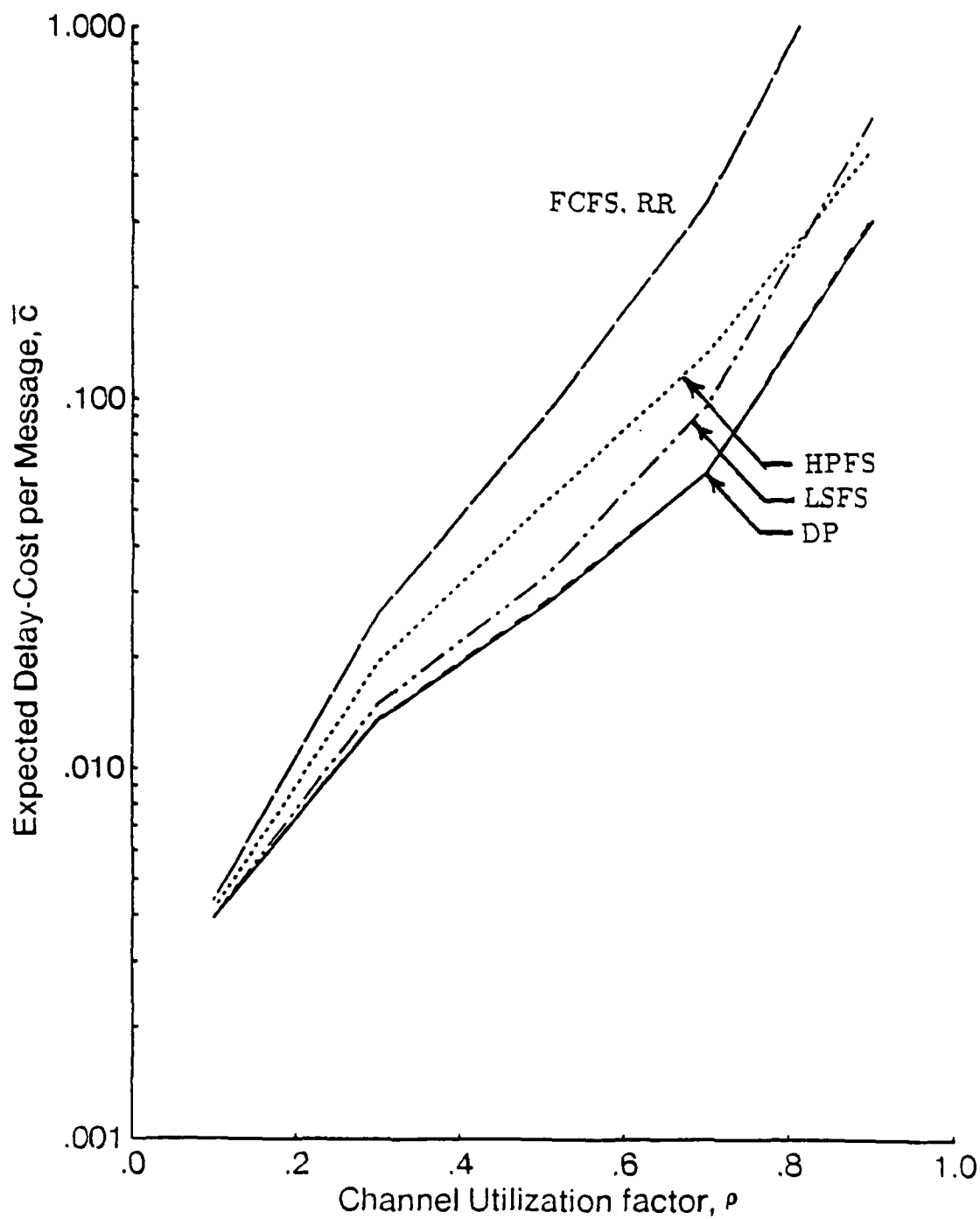


Fig. 2.10 Expected cost per message versus load factor ρ for fixed-length messages, ramp cost functions, $\bar{C}_a = 5\bar{T}$, $\bar{L} = K$, and $\bar{C}_p = 1$.

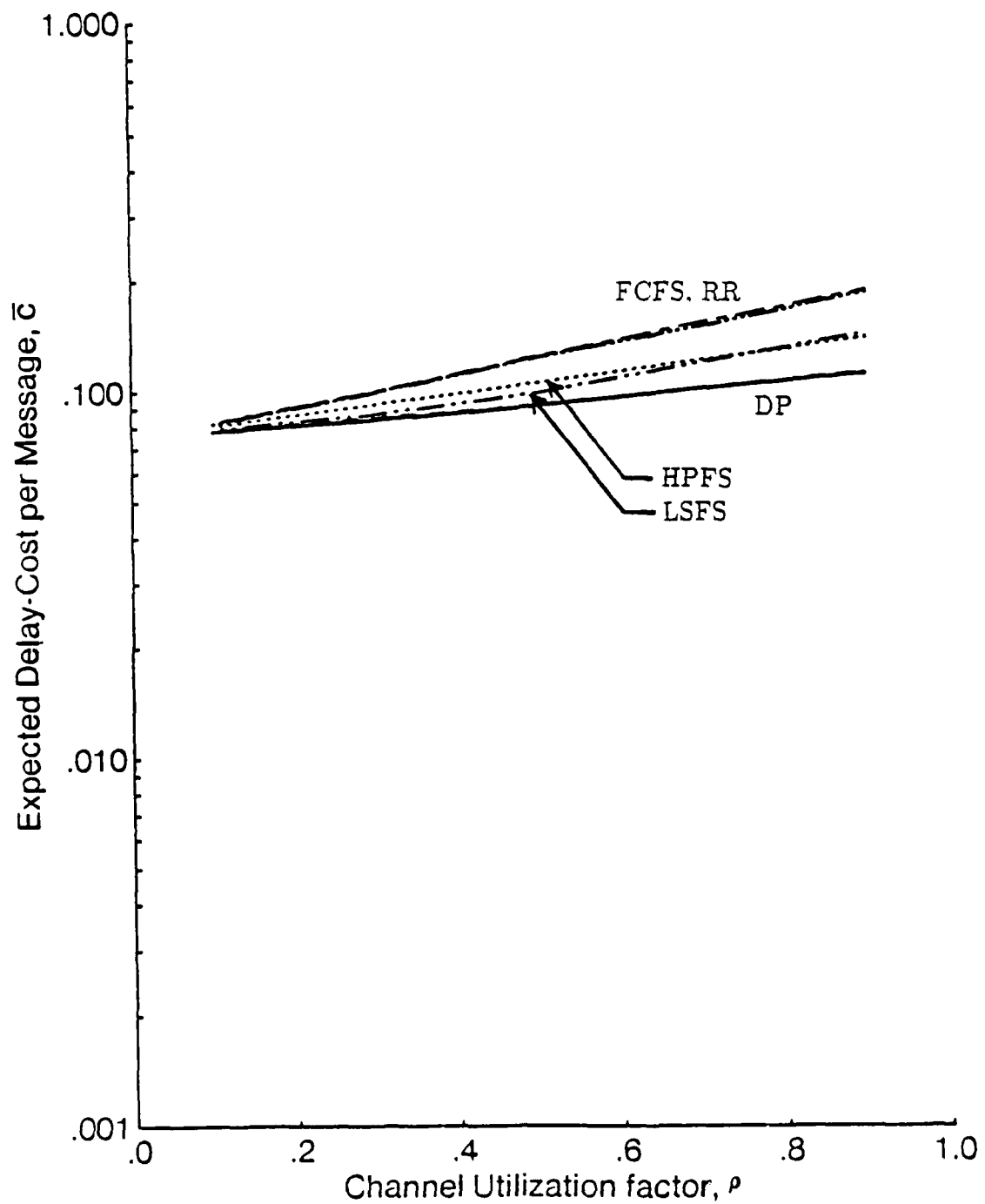


Fig. 2.11 Expected cost per message versus load factor ρ for exponentially-distributed message lengths, step cost functions, $\bar{C}_a = 5\bar{T}$, $\bar{L} = K$, and $\bar{C}_p = 1$.

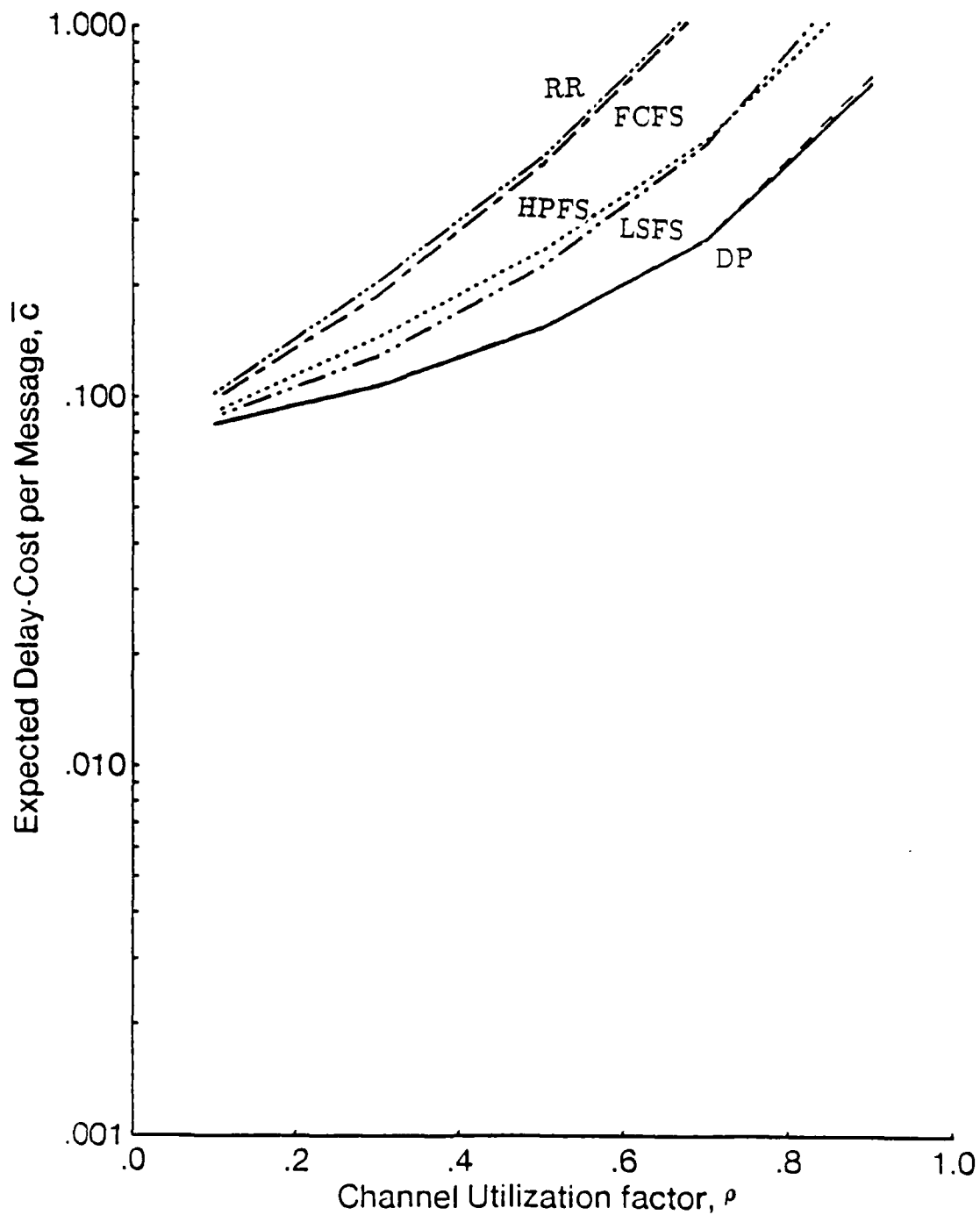


Fig. 2.12 Expected cost per message versus load factor ρ for exponentially-distributed messages, ramp cost functions, $\bar{C}_a = 5\bar{T}$, $\bar{L} = K$, and $\bar{C}_p = 1$.

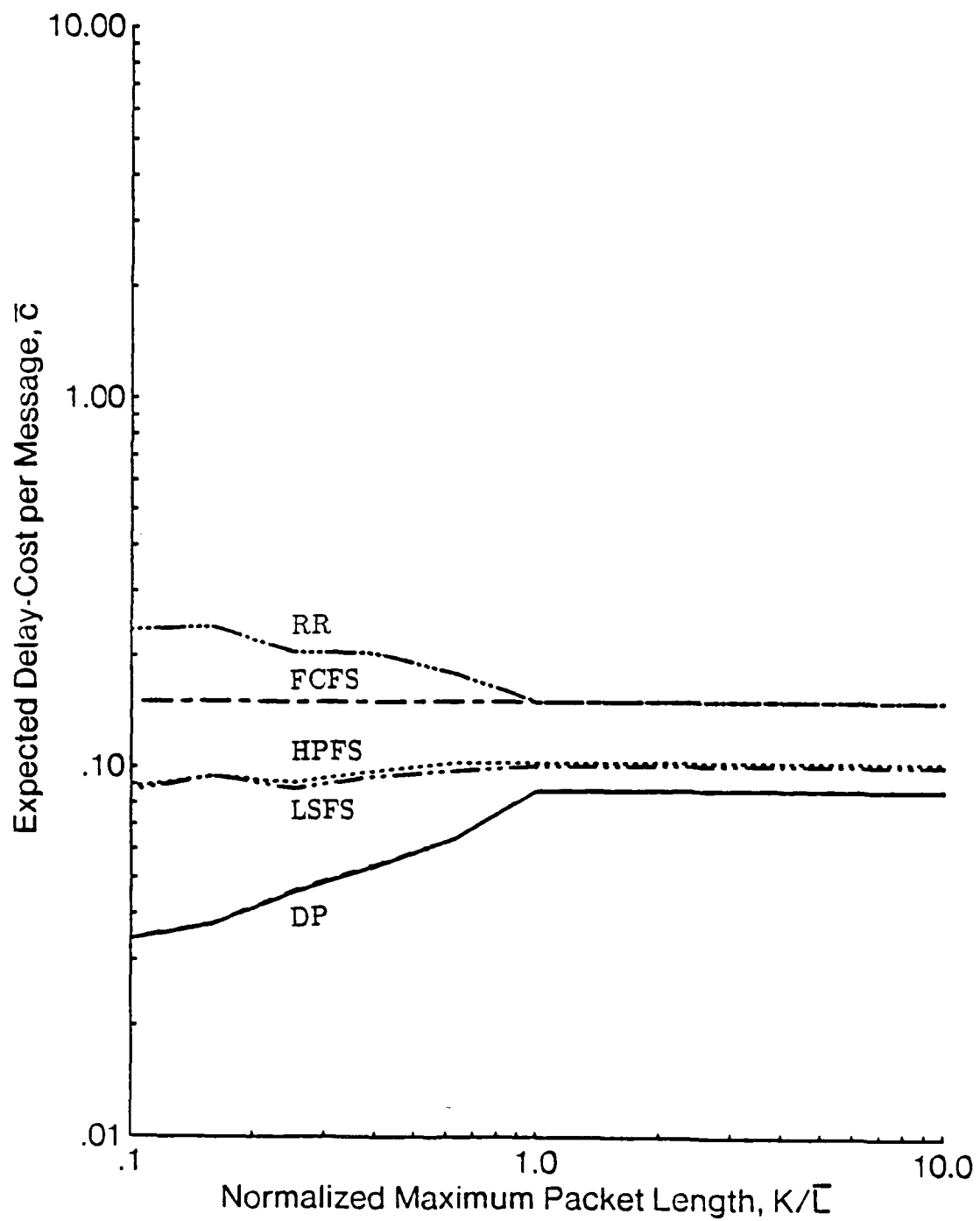


Fig. 2.13 Expected cost per message versus normalized maximum packet length $(k/L) \rho$ for fixed-length messages, step cost functions. $\rho = 0.3$, $\bar{C}_a = 5\bar{T}$, and $\bar{C}_p = 1$.

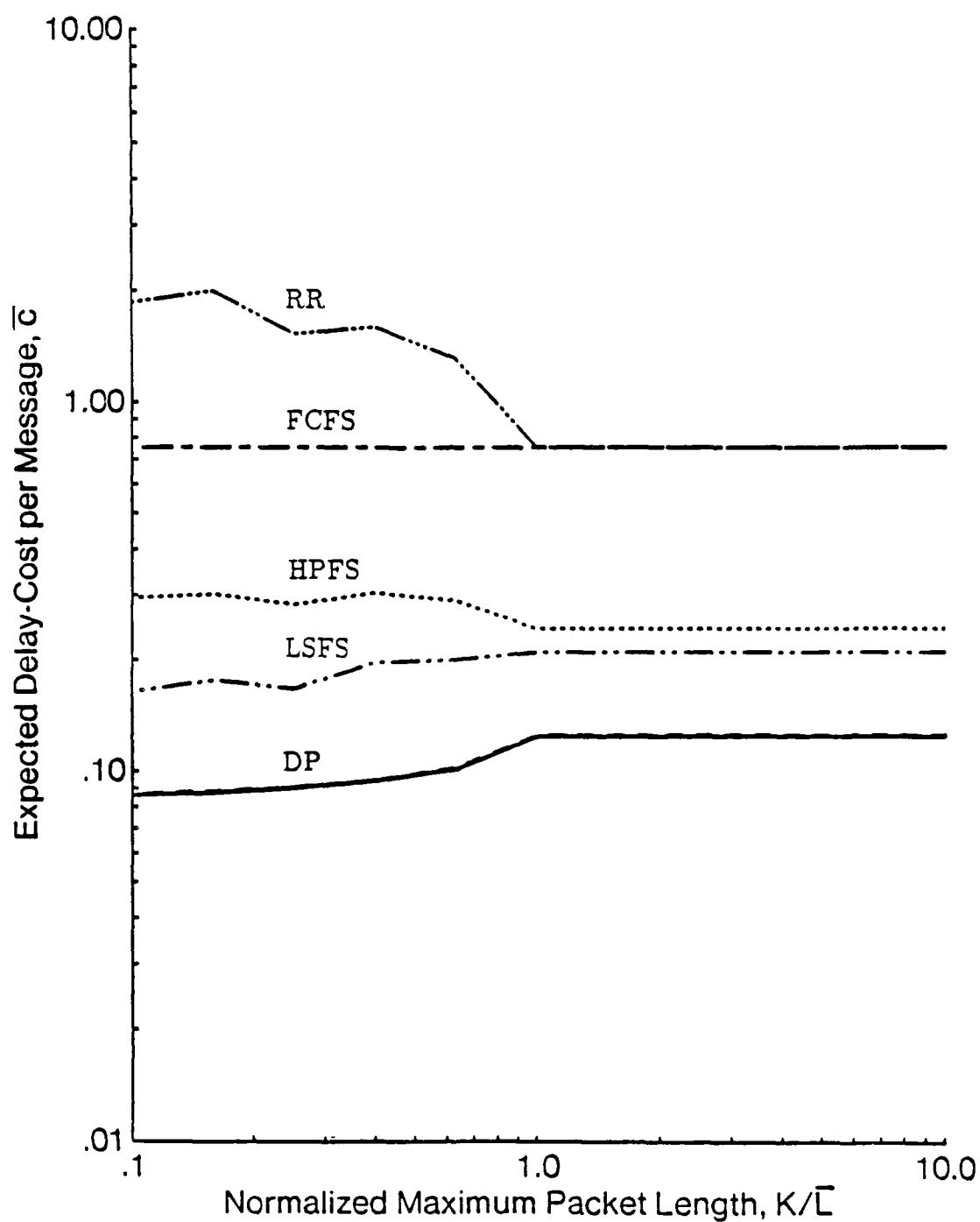


Fig. 2.14 Expected cost per message versus normalized maximum packet length ($K/\bar{L}$) for fixed-length messages, ramp cost functions, $\rho = 0.3$, $\bar{C}_a = 5\bar{T}$, and $\bar{C}_p = 1$.

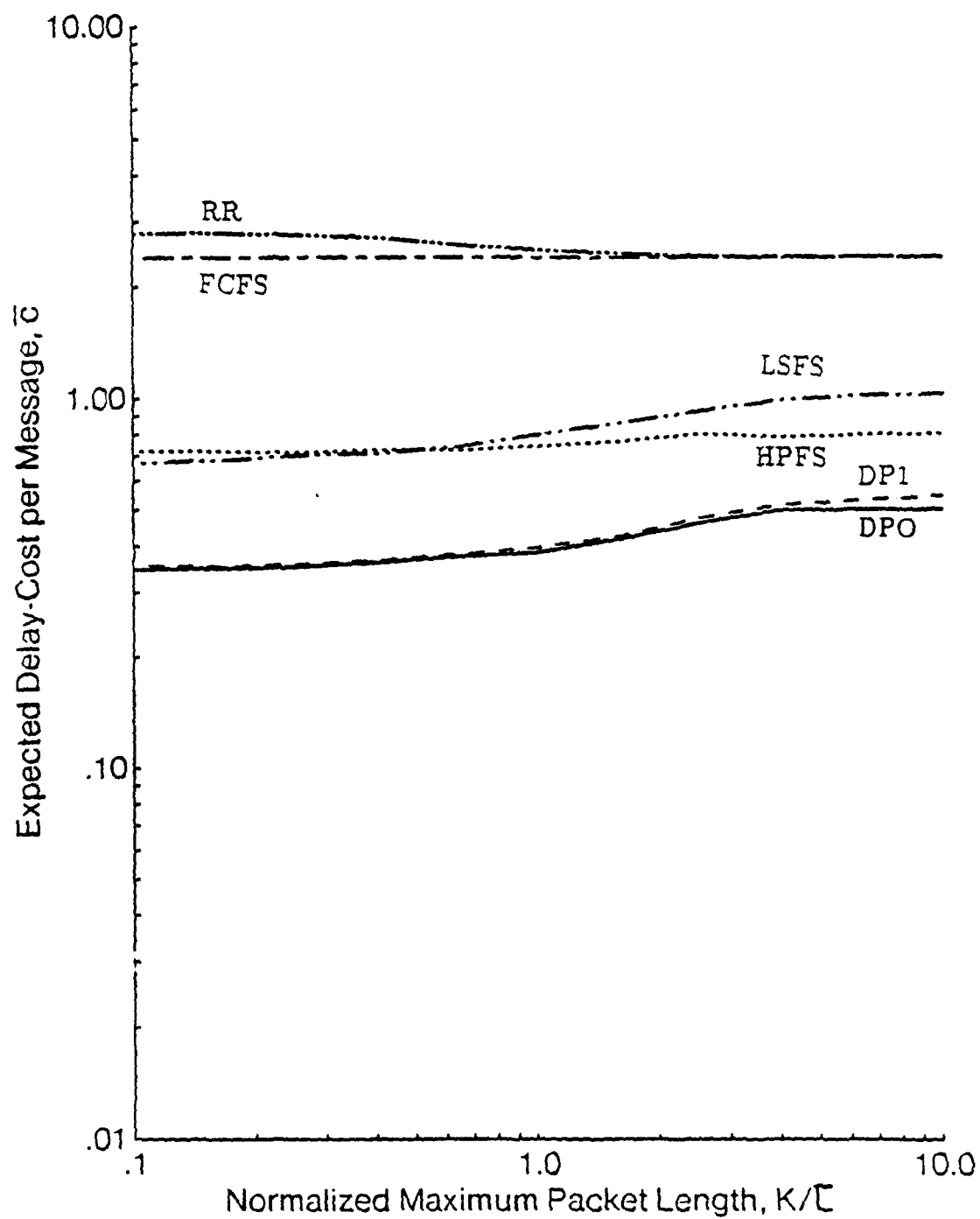


Fig. 2.15 Expected cost per message versus normalized maximum packet length ($K/\bar{L}$) for exponentially-distributed message lengths, step cost functions, $\rho = 0.8$, $\bar{C}_a = 5\bar{T}$, and $\bar{C}_p = 1$.

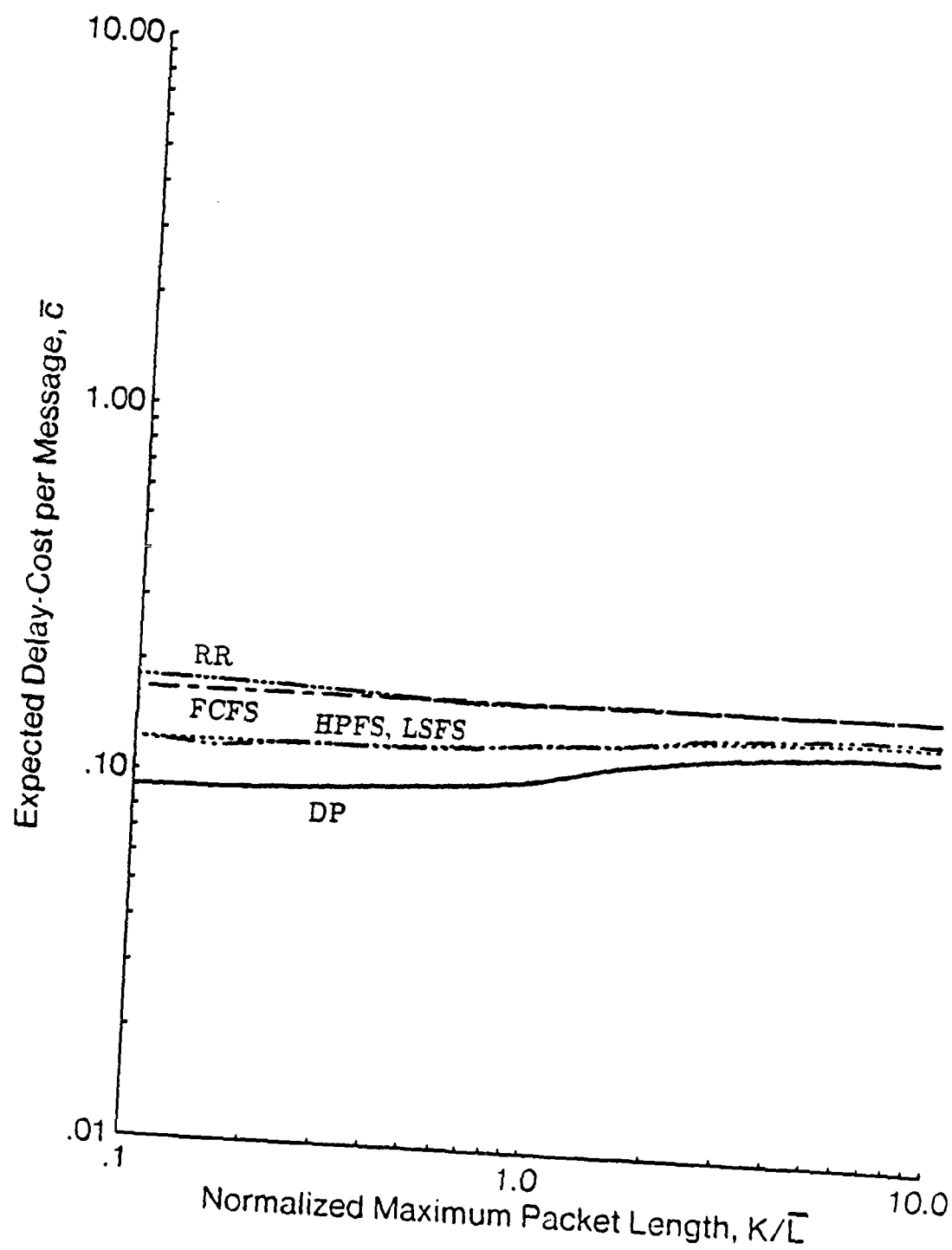


Fig. 2.16 Expected cost per message versus normalized maximum packet length ($K/\bar{L}$) for exponentially-distributed messages, ramp cost functions, $\rho = 0.8$, $\bar{C}_a = 5\bar{T}$, and $\bar{C}_p = 1$.

with decreasing values of $K/\bar{L}$. For all other policies, there is a tradeoff between the gain achieved with a large number of decision opportunities and the loss incurred with the increase of mean message delay. The DP policies show consistent improvement which gets more significant as $K/\bar{L}$ decreases. The change in performance achieved in LSFS and HPFS is not as certain, and HPFS in the case of a ramp cost function actually degrades with decreasing values of $K/\bar{L}$ (figures 2.13 and 2.16). The results obtained for exponential message lengths are similar to those discussed above, except that $\bar{c}$ is not constant for $K/\bar{L} > 1$.

2.4 Summary

The problem of scheduling of message transmissions in broadcast local area networks was considered here. By assigning cost to message delays, we formulated a transmission scheduling problem. The objective was to find a scheduling policy which minimized the expected cost per message. Finding such an optimal scheduling policy is analytically intractable. We proposed two suboptimal scheduling policies. The first is based on the optimal static scheduling and can be used for fixed-length messages and for step cost functions. The second policy is a generalization of the $C\mu$ rule and can be used for general message-length distributions and general cost functions. Using regenerative simulation, we showed that both policies outperform other common scheduling policies. In the next chapter we show that how these policies can be implemented in a distributed environment.

Chapter 3

Broadcast Local Area Networks incorporating Scheduling Policies

In the previous chapter we studied a number of scheduling policies under ideal conditions (i.e., perfect and instantaneous knowledge of relevant information: messages, their cost, penalty, etc.). In this chapter we are concerned with the incorporation of these policies in broadcast local area networks in view of integrating services: i.e., supporting different applications and satisfying their requirements. There are two issues involved. One is the provision of efficient channel access schemes for broadcast local area networks operating at high speed. The other issue is the provision of necessary mechanisms for the implementation of scheduling policies. There are efficient channel access schemes based on the attempt-and-defer access mechanism which solve the first problem and which constitute the basis for our solution. With respect to the second problem, we choose to consider a distributed implementation for two reasons. The first is reliability: with distributed scheduling, the system is not dependent on the proper operation of a central scheduler. The second is improved performance, in terms of both delay and channel utilization. This is due to the fact that, in distributed scheduling, only stations' local information must

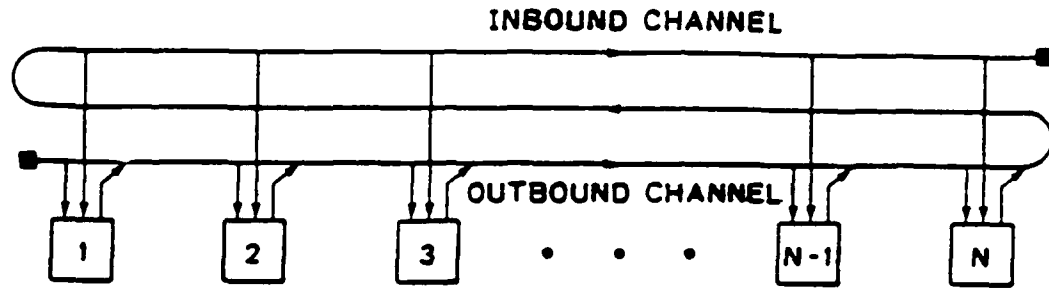
be broadcast on the network to the central scheduler, and the control information from the central scheduler back to the stations is not required. We propose in this chapter a scheme which includes (i) a mechanism for the dissemination of all relevant information required by the scheduling policy, and (ii) a station architecture which allows the implementation of any of the scheduling policies described in the previous chapter.

Before we present the distributed description of the scheme, we begin by giving a short review of the attempt-and-defer access schemes for high speed LAN. Then, we give a brief discussion of the basic requirements for distributed implementation of scheduling policies.

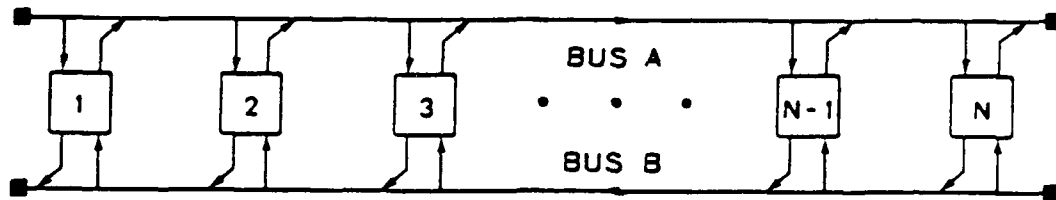
3.1 Attempt-and-Defer Access Schemes

High-performance *bus-oriented* local area networks have recently become feasible due to the emergence of a new class of demand assignment multiple-access (DAMA) schemes [FINE84] which, unlike CSMA/CD [STAL84, METC76] and the IEEE token-passing schemes [STAL84, IEEE83] can effectively utilize high bandwidths.

The new DAMA schemes are based on the attempt-and-defer mechanism [FINE84] and employ a *unidirectional bus structure*. Broadcast communications is achieved in two ways. The first approach is to fold a unidirectional cable onto itself so as to create two channels, an *outbound channel* onto which the users transmit packets and an *inbound channel* from which users receive packets; all signals transmitted on the outbound channel are returned on the inbound channel, as shown in Fig. 3.1(a). This will be referred to as the *folded bus* structure. Another way to achieve broadcast communications is to provide two unidirectional busses with signals propagating in



(a) Folded Bus Structure



(b) Dual Bus Structure

Fig. 3.1 Unidirectional bus structures.

opposite directions as shown in Fig. 3.1(b). This will be referred to as the *dual bus* structure.

In both the folded bus and dual bus structures, on every unidirectional bus segment to which stations are connected, there exists an implicit ordering among the stations which the DAMA schemes under consideration make use of. These schemes also require that on every unidirectional bus segment, each station be given

the ability to sense activity due to stations on the upstream side of its transmit tap. In the dual bus structure the receive taps provide this function while in the folded bus structure, as shown in Fig. 3.1(a), additional sense taps are needed.

The attempt-and-defer access mechanism basically operates as follows. A station wishing to transmit on a given bus waits until that bus is idle. It then begins to transmit, thus establishing its desire to acquire the bus. However, if another transmission from upstream is detected, the station aborts its transmission and defers to the one from upstream. The most upstream transmission is therefore allowed to continue conflict-free. Examples of networks using the attempt-and-defer access mechanism are Expressnet [TOBA83], Fasnet [LIMB82], U-Net [GERLS3a], Token-Less Protocols (TLP) [RODR84], MAP [MARS82], and Buzznet [GERLS3b].

As an example, consider the access scheme of Expressnet. The synchronizing event used by stations to determine when to transmit is the end-of-carrier on the outbound channel (EOC(out)). All stations with packets ready to be transmitted attempt to transmit when they detect EOC(out) and defer to upstream transmissions. Clearly, a station that has completed the transmission of a packet in a given round will not encounter the event EOC(out) again in that round, thus guaranteeing that no station will transmit more than once in a given round. Define a *train* to be the succession of transmissions in a given round. Due to the special bus arrangement shown in Fig. 3.1(a), the end of a train on the inbound channel (EOT(in)) will visit each station in the same order as they are permitted to transmit. To start a new round, EOT(in) is used as the synchronizing event, in the same way that EOC(out) was used in the above description. Thus stations synchronize their transmissions to the first of two events, EOC(out) or EOT(in). (Note that, at a given station, only one such event can occur at a time.)

Besides Expressnet, MAP also uses a folded bus structure; all other systems use

the dual bus structure. U-Net, TLP, and Buzznet require that packets be transmitted simultaneously on both busses; at any one time, however, the access arbitration is performed only on one of the busses. In Fasnet, the busses are accessed independently; the transmissions are time-slotted; on each bus, the most downstream station, by setting a bit at the beginning of a slot on the other bus, indicates the end of round to the most upstream station which in turn starts a new round.

The most significant difference among attempt-and-defer protocols is in the way new rounds are generated [FINE84]. Otherwise, using any of these protocols, packet transmissions in a round are without overlap and with a minimum gap separating them. A gap equal to the round trip delay is incurred only once at the end of a round. This is in contrast with the CSMA/CD or IEEE 802 token passing scheme, in which gaps of the order of the round trip delay occur between any two consecutive packet transmissions.

3.2 Basic Requirements for the Distributed Implementation of Scheduling Policies

In order to implement a scheduling policy in a distributed environment, (e.g. a broadcast LAN such as those described above) one must provide two additional basic functions; namely: (i) a broadcast mechanism for the dissemination of relevant information concerning the traffic to be transmitted, and (ii) the scheduling logic to be implemented in each station.

A. Broadcast Mechanism

In order to be able to transmit a message, a station must first broadcast information to all other stations regarding this message. This information includes the generation time, the cost function, the length of the message, and the identity

of the source station. For all practical purposes, we may refer to the transmission of such information as a request. It is important that the mechanism devised be simple, robust, reliable, and efficient.

As requests are generated by the various stations randomly, the transmission of such requests invoke the use of channel access protocols that (i) must be efficient and (ii) do not introduce excessive delays. From a practical point of view, it is most likely that all traffic will belong to a certain finite number of classes, whereby the cost function for all messages in a given class is the same and can be considered known a priori to all stations. This allows the cost function of a message to be replaced by its class.

Furthermore, it is clear that if the information regarding requests is made more implicit, the broadcast mechanism becomes more robust. For example, as we shall see below when presenting our proposal, stations do not require the complete knowledge about the identity of source stations. Every station just needs to know if a received request belongs to it or some other station. This information can easily be obtained if the station knows the delay from its transmitter to its receiver, and stores the times at which it has made requests. As another example, the generation time of a message may be deduced (to a sufficient accuracy) from the time at which a request is received (or transmitted), therefore obviating the need to transmit such information explicitly and to keep all stations clock synchronized.

B. Stations' Scheduling Logic

Clearly, for a distributed implementation of scheduling policies, it is necessary for the scheduling algorithm to be implemented in each station. This requires special consideration concerning the stations' architecture. First of all, we require the architecture to be general enough to accommodate a variety of scheduling policies. Secondly, the architecture and its components must be simple enough to allow VLSI

implementation thereof. Finally, the architecture and its implementation should be efficient in that the time it takes the scheduling algorithm to reach the result be short enough so as not to affect the overall performance.

For the distributed implementation of a scheduling policy based on the above two functions to operate properly, a requirement referred to here as *consistency* must be satisfied. This means that the result of the scheduling algorithm within a station as to which message is to be transmitted at a given point in time (relative to that station's location) must be consistent with the same at all other stations. This concept is further elaborated upon when we describe our design.

3.3 A Design Meeting The Requirements

In this section we describe a proposal for a high-speed network incorporating the transmission scheduling policies. The basis of this proposal is Expressnet with its doubly folded bus topology. We consider the various applications to belong to a number of classes. All messages in a class have the same cost function considered to be known a priori to all stations. Furthermore, to simplify the design, we consider a slotted system (similar to Fasnet) and assume that all messages consist of a single packet whose transmission time is equal to the slot size.

3.3.1 Broadcast Mechanism

Consider the folded bus configuration shown in Fig. 3.1(a). We assume that upstream to each transmitter there is a receiver. The most upstream user transmits a clock signal which keeps the system bit synchronous. From this clocking information, users listening to the channel are able to identify fixed length slots traveling

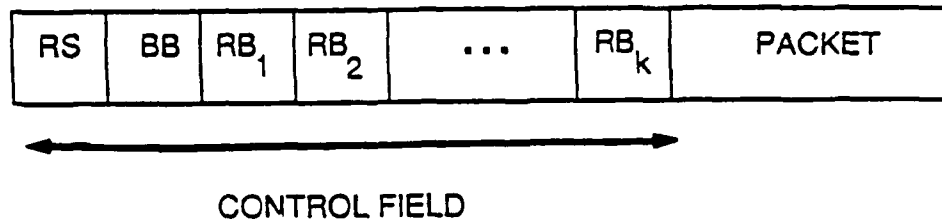


Fig. 3.2 Slot Format.

downstream. Each slot begins with a control field. (See Fig. 3.2.) The control field consists of a busy bit (BB), reset bit (RS), and one request bit for each of the traffic classes supported by the network. We denote the request bit of class- k traffic by RB_k .

Two algorithms for setting the request bits are proposed resulting in two service disciplines: most-upstream-first-served (MUFS) and round robin. In the MUFS approach, a station with a backlogged class- k request sets and reads RB_k . If it reads $RB_k=0$, then it was successful. Otherwise, it should try again in the next slot.

In the round-robin version the access protocols corresponding to different classes run independently. We describe the access protocol corresponding to class k . We first describe how the stations access the class- k request bits within the round and then describe the mechanism by which a new round is initiated. Within a round a station can set RB_k of a slot, only if in the previous slot it reads $RB_k = 1$. In that case it simultaneously sets and reads RB_k ; setting an already set bit is assumed to have no effect. If it reads $RB_k = 0$, then it was successful. Note that within a round, every station may set RB_k only once.

The mechanism for initiating a new round is similar to Expressnet. Define a train to be a succession of $RB_k=1$ in a given round. The end of a train on

the inbound channel. $EOT(in)$, is detected whenever an $RB_k=1$ is followed by an $RB_k=0$. To start a new round, $EOT(in)$ is used as the synchronizing event. To avoid losing synchronization, after detecting $EOT(in)$, all stations set the first RB_k on the outbound channel. Therefore the first RB_k in such round is ignored. Note that at the beginning of the round, although the most-upstream station does not read $RB_k(out)=0$, it nevertheless accesses the request bit as if it was set.

3.3.2 Station Architecture

The station architecture is depicted in Fig. 3.3. For each class of traffic there exists a local request queue and a shift register. In addition to these, there is a packet buffer, a global request queue, and a scheduler. When a packet is generated, it is put in the packet buffer. At the same time, a request is put in the local request queue of the appropriate class. The request which is put in the local queue is the pointer to the packet in the packet buffer. Requests are broadcast using one of the mechanism described above. Recall that broadcasting a request of class k corresponds to setting the class- k request bit RB_k . When the station successfully sets RB_k , it removes the request from the head of class- k local queue and inserts it in the class- k shift register. At all other times a nil pointer is inserted in the shift registers.

The delay through each of the shift registers is equal to the propagation delay from the transmitter of the station to its receiver. (Note that, in the folded bus topology of Fig. 3.1(a), this delay is the same for all stations.) Whenever a set RB_k is observed on the inbound channel, the output of the class- k shift register is removed and tagged with the broadcast-time of the request and its class, and put in the global queue. ($RB_k=1$ at the beginning of the round under the round robin algorithm is ignored.) Note that if the pointer is nil the request belongs to

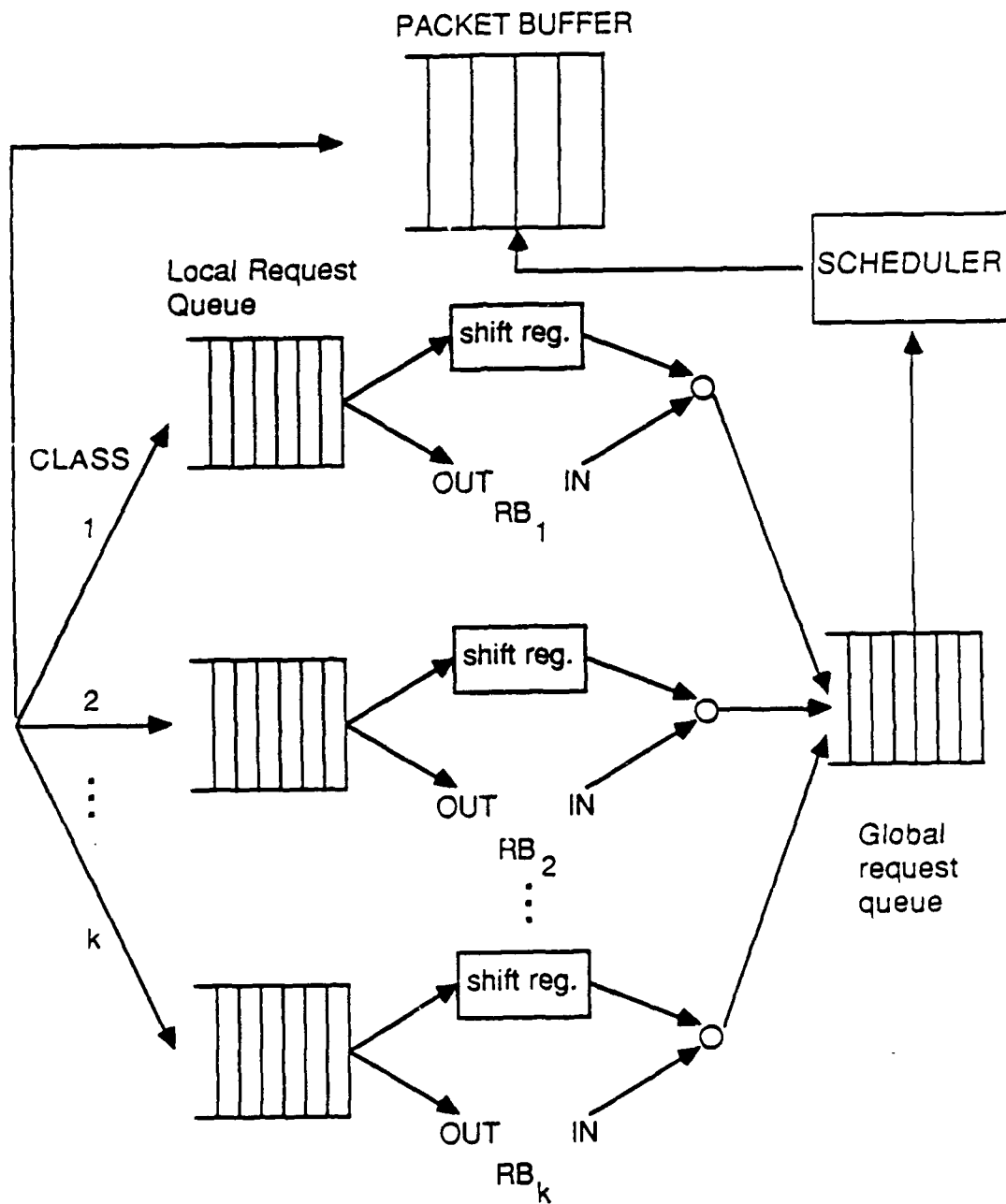


Fig. 3.3 Station architecture

another station, otherwise it belongs to the same station. (Hence, there is no need to transmit station identification in the request.)

The broadcast-time of the request, which is equal to the time the request is received on the inbound channel minus the round-trip delay, is used as an estimate for the generation time of the packet. The inaccuracy in this estimate is equal to the queueing delay at the local queues. The queueing delay at every local queue is a function of the load factor at that queue, i.e., the average number of requests generated in that class. Clearly, the lower the load factor is, the smaller the queueing delay is. If the total traffic is uniformly distributed among classes, the load factor for each of the classes is low. If, on the other hand, a large fraction of traffic is in one class, then the load factor for that class is large. To reduce the load, more than one request bit may be allocated to that class; this has the effect of dividing the load into several queues, each with lower load and delay.

Each request in the global queue consists of a packet pointer (possibly nil), a class, and an estimate of generation time. The class specifies the cost function of the message. The scheduler uses the cost functions and the (estimated) generation times to select a packet for transmission in the next slot. If the pointer of the scheduled packet is nil, then the scheduled packet belongs to the some other station. Otherwise, the packet belongs to the same station and must be transmitted in the next slot.

Due to the special feature of the folded bus topology, there is consistency among the stations. To illustrate this, we consider the stations S1 and S2, where S1 is upstream to S2. The propagation delay from the transmitter of S1 to the transmitter of S2 is equal to the propagation delay from the receiver of S1 to the receiver of S2. Recall that the content of the global queue at a station is only a function of the requests received by that station on the inbound channel. Therefore, the state of

the global queue at S1 when a slot arrives at S1, is identical to the state of global queue at S2 when the same slot reaches S2.

Before transmitting the packet in the slot, the station checks the busy bit (BB) of the slot. If BB is not set, the station sets it and transmits its packet. If BB is already set, it means that the slot is used by some upstream station. Such a conflict can occur only if there has been inconsistency among the global queues. In that case, the station detecting the error sets the reset bit (RS), so that the head station broadcast the content of its global queue. The RS is also used by the stations which are becoming active to acquire the content of the global queue. Thus, if inconsistency is to occur due to errors, the redundancy created by BB allows the stations to detect it and recover from it.

3.4 Summary

In this chapter, we presented a broadcast local area network for integrated-services applications. This network was based on the distributed implementation of scheduling policies. The proposed scheme included (i) a mechanism for the dissemination of all relevant information required by the scheduling policy, and (ii) a station architecture which allowed the implementation of the scheduling policies described in the previous chapter. The broadcast protocol was based on the attempt-and-defer access mechanism. In the next chapter, we will present a number of fiber optic configurations for local area networks, in particular those based on the attempt-and-defer access mechanism.

Chapter 4

Fiber Optic Configurations for Local Area Networks

The medium that provides high bandwidth in local area networks is naturally fiber optics. While this technology lends itself easily to networks with point-to-point links, (e.g., rings) with broadcast networks the problem is more complex. Due to intrinsic characteristics of certain fiber-optic components, in particular the reciprocity of couplers and the low impedance of detectors, the use of this technology in the implementation of a shared multitapped bus requires careful design consideration. It has been known that, due to the reciprocity and excess loss of optical couplers, the number of stations that can be accommodated on a linear fiber optic bus is severely limited [HUDS74, MILT76, ALTM77, AURA77, VILL81, SCHM83, LIMB84]. In this chapter, we first examine the unidirectional networks in question (folded bus and dual bus topologies) and extract features which prove useful in devising fiber optics configurations. We then present various fiber optics configurations for the implementation of the unidirectional broadcast networks. In the following section we provide a complete power budget analysis of fiber optic configurations. This leads to the performance evaluation of the configurations in terms of the maximum

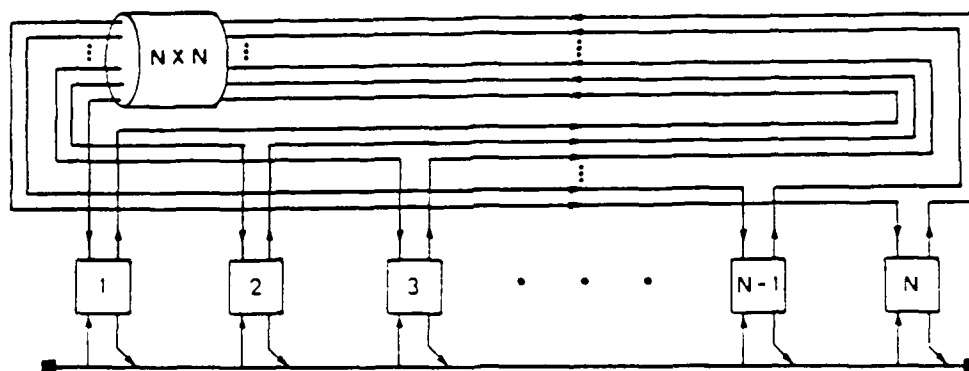


Fig. 4.1 Star configuration.

number of stations that can be supported for a given power margin.

4.1 Feature Abstraction

In the folded bus structure shown in Fig. 3.1(a), we distinguish two basic functions. The first is that of providing connectivity among all stations for *data transmission*; that is, the provision of data transmission paths among the stations' transmitters and their receivers for broadcast communications. The second function is that of providing the required feature for *access control*; this consists of achieving a physical order among the stations according to which the latter use the synchronizing events.

In the particular configuration shown in Fig. 3.1(a), the data transmission and access control functions are combined into a single linear bus, with appropriately placed sense taps, transmit taps and receive taps. However, these two functions may be implemented in two separate networks of possibly different topologies: a *data subnetwork* and a *control subnetwork*. For example as shown in Fig. 4.1, the

first may have a star topology satisfying the full connectivity requirement, and the second may be a linear bus achieving the linear order requirement. (The star coupler shown in the figure simply distributes the power from its input arms into its output arms.) In the next section, we consider a number of fiber optic configurations for the control and data subnetworks.

The dual-bus structure of Fig. 3.1(b) can be viewed as composed of two linear control subnetworks (as defined for the folded bus structure) and thus can use the control configurations proposed in section 2. Also, the data configurations proposed in that section can be employed independently for access schemes which do not require any linear ordering among stations. Examples of such schemes are PODA [JACO78] and HAM [HANS81], both of which achieve a high utilization of bandwidth through a reservation technique.

It is known that the throughput performance of an attempt-and-defer scheme degrades as the time t_d that a station takes to detect EOC(out) on the control subnetwork increases [TOBA83]. Since the bandwidth needed in the control subnetwork is lower for larger values of t_d , it is possible in the folded bus structure to trade throughput for (lower) bandwidth of the control subnetwork, which may then be implemented using cheaper optical transmitters and receivers, or a medium other than fiber optics such as twisted pair or coaxial cable. (If the control subnetwork is implemented in fiber optics, at every station, instead of using separate transmitters for the data and control subnetworks, we may use only one transmitter and split its output through a coupler between the two subnetworks.)

In all data configurations that we consider, as is also the case in the original structure of Fig. 3.1(a), the data transmission subnetwork can be separated into two parts: the *data collection subnetwork* and the *data distribution subnetwork*. This partitioning of the data transmission network implies that there exists a single

funneling point separating the transmitters from the receivers, through which all packet transmissions flow, and such that the collection subnetwork connects all transmitters to this point, and the distribution subnetwork connects this point to all receivers. For example, for the linear bus structure, the funneling point is the bus segment connecting the outbound channel to the inbound channel, and in the star configuration, the star coupler constitutes the funneling point.

Regardless of the particular network structure implementing the folded-bus structure, to maintain its performance characteristics, it is required that all packets in a round flow through the funneling point without overlap and with a minimum gap separating them. While this requirement is naturally achieved in the structure shown in Fig. 3.1(a), for a structure which separates the data transmission network from the control network, it implies the following: for any pair of stations S_i , S_j , $j > i$, the propagation delay from S_i to the funneling point on the data collection subnetwork must be equal to the sum of the propagation delay from S_i to S_j on the control network and the propagation delay from S_j to the funneling point on the data collection subnetwork. If the data and control subnetworks are both implemented using fiber optics, this requirement translates into a set of equalities in terms of fiber lengths, and can easily be satisfied by having the control and data collection fibers in the same cable.

In the case of Expressnet, there is an additional constraint: in order to implement the start of round procedure based on EOT(in) as was described in chapter 3, section 3.1, the delay through the data transmission subnetwork from a station back to itself must be the same for all stations. This requirement is clearly satisfied by the star configuration shown in Fig. 4.1.

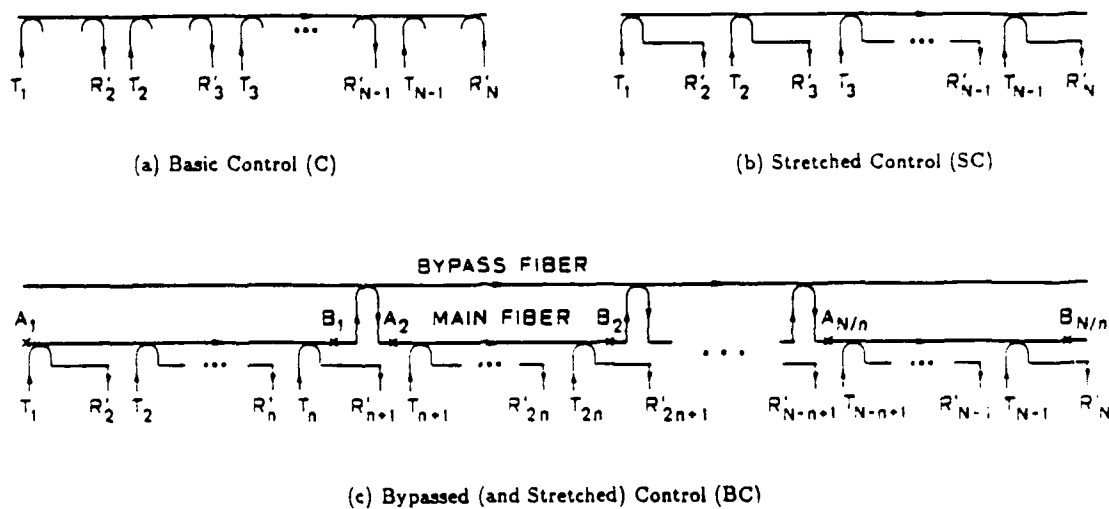


Fig. 4.2 Configurations for the control subnetwork.

4.2 Fiber Optic Configurations

4.2.1 Configurations for the Control Subnetwork

A direct implementation of the control bus (as shown in the star configuration of Fig. 4.1), is depicted in Fig. 4.2(a) and is referred to as the *basic control* (C) configuration. Due to excess loss and reciprocity of couplers, the maximum number of stations that can be supported by this configuration is quite limited. Here, we propose two techniques to overcome this limitation and also discuss the use of non-reciprocal couplers.

4.2.1.1 Stretched Configuration

The number of couplers needed to implement the control subnetwork may be decreased by a factor of two from that of the basic control configuration by considering the configuration shown in Fig. 4.2(b). In this configuration a single coupler is

used per station (for its transmitter), with the output of its otherwise unused port connected via an additional *stretch* of fiber to the sensor of the next downstream station. (The sensor of a station must be able to sense the activity due to only the upstream users.) This configuration is referred to as the *stretched control* (SC) configuration. In a typical local area network, the cost of the additional fibers is outweighed by the saving in the number of couplers.

4.2.1.2 Bypassed Configuration

The number of couplers between transmitters and receivers may be reduced by providing bypass paths as shown in Fig. 4.2(c). This control configuration is referred to as the *bypassed (and stretched) control* (BC) configuration. Assume that there is an equal number of couplers, n , on each bypassed segment. Then the number of couplers between T_1 and R'_N over the signal path with the least number of couplers is $2n + \frac{N-1}{n} - 1$, which is minimized at $n = \sqrt{\frac{N-1}{2}}$. Therefore, the number of couplers between any transmitter-receiver pair is less than or equal to $2\sqrt{2(N-1)} - 1$, as compared to $N - 1$ in the SC configuration. (The couplers on the bypass bus may themselves be bypassed so as the bound on the number of couplers between any transmitter-receiver pair be of the order of $\sqrt[3]{N}$, and so on.) bypassing improves the performance of the configuration, for sufficiently large values of N .

In the bypassed control configuration, the variation in propagation delay over all signal paths between every transmitter and receiver pair must be much less than a bit period. Otherwise, the resulting ambiguity in the bit timing may significantly increase the bit-error rate. We note however that, even at high data rates, the propagation delays can be sufficiently equalized by simply having the *bypass* and *main fibers* in the same cable. (Indeed we observe that, when the data rate is 200 Mbps for example, one bit period corresponds to the propagation delay over 1 meter

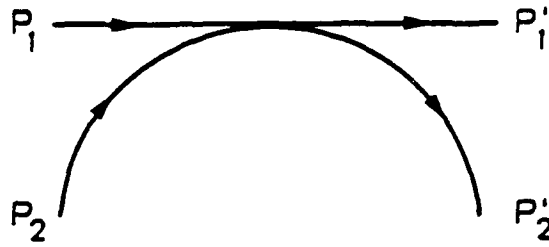


Fig. 4.3 Optical coupler.

of fiber.) Furthermore, we cannot have both bypass and main fibers be singlemode when laser sources are used. Otherwise, the temporal coherence between the signals received over separate optical paths can result in fading at the receivers.

4.2.1.3 Non-reciprocal Couplers

Besides excess loss, the factor which limits the number of stations on a linear bus is the reciprocity of the couplers. Reciprocity demands that the fraction of power removed from a fiber be equal to the fraction of power injected into that fiber. This is clearly not desirable, since more stations can be supported if every station injects most of its power into the bus even if it needs to remove a small fraction of the power from the bus. With the notation defined in Fig. 4.3, the input-output relationship of a reciprocal (or symmetric) coupler is given by

$$\begin{bmatrix} P_1' \\ P_2' \end{bmatrix} = \begin{bmatrix} \alpha(1-x) & \alpha x \\ \alpha x & \alpha(1-x) \end{bmatrix} \begin{bmatrix} P_1 \\ P_2 \end{bmatrix} \quad (4.1)$$

where α denotes the transmittance due to an excess loss of E_C dB: i.e. $\alpha \triangleq 10^{-E_C/10}$.

Recently a number of non-reciprocal (or asymmetric) couplers have been designed and implemented. In one such implementation, the light is strongly coupled

from a singlemode fiber to only a few individual modes of a multimode fiber. However, only the small fraction of the power propagating in these individual modes is coupled out of the multimode fiber [WOOD85]. In another implementation, an asymmetric biconical taper coupler is made by using fibers of different core diameters or by tapering the fibers by different amounts [RIVE84]. Any coupling matrix which cannot be put into the form of the matrix in (4.1), corresponds to an asymmetric coupler (as long as the columns of the matrix add up to less than one). For the sake of simplicity we only consider the asymmetric couplers which have the following characteristics:

$$\begin{bmatrix} P'_1 \\ P'_2 \end{bmatrix} = \begin{bmatrix} \alpha x & \alpha x \\ \alpha(1-x) & \alpha(1-x) \end{bmatrix} \begin{bmatrix} P_1 \\ P_2 \end{bmatrix} \quad (4.2)$$

The error made in approximating the true characteristic of asymmetric couplers by (4.2), can be bounded by augmenting the excess loss.

4.2.2 Configurations for the Data Subnetworks

Here, we study a number of fiber optic configurations for the data subnetworks. Based on their topologies, the configurations considered here can be grouped into three classes: *basic*, *hybrid*, and *compound*. In a basic configuration the same topology is used for both the collection and distribution subnetworks. In a hybrid configuration, different topologies are used for the collection and distribution subnetworks. In a compound configuration a set of collection subnetworks is connected, via a star, to a set of distribution subnetworks.

4.2.2.1 Basic Configurations

The first basic configuration we consider is the *linear* (L) configuration, in which both the collection and distribution subnetworks have a linear-bus topology.

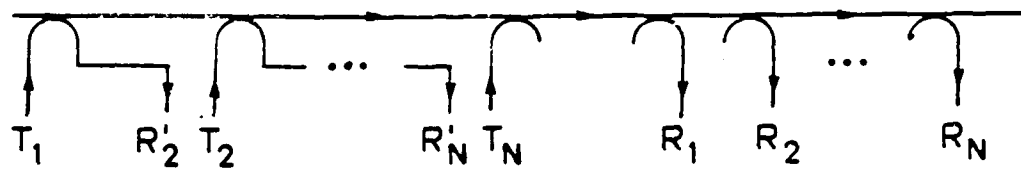
with the control function provided by the (linear) collection subnetwork. This configuration is shown in Fig. 4.4(a). Although not shown in Fig. 4.4(a), one may use the bypass technique proposed in the previous section to implement the collection-control subnetwork. In the configurations considered in this section only the stretch feature is used; the effect of the bypass feature on the performance is well understood from the previous section.

The second basic configuration to implement a data subnetwork is the *star* (S) configuration, as shown in Fig. 4.4(b). For the sake of clarity, the fiber lengths in the data configurations shown in this section are not adjusted to achieve the same delay from every station back to itself, as required by Expressnet.

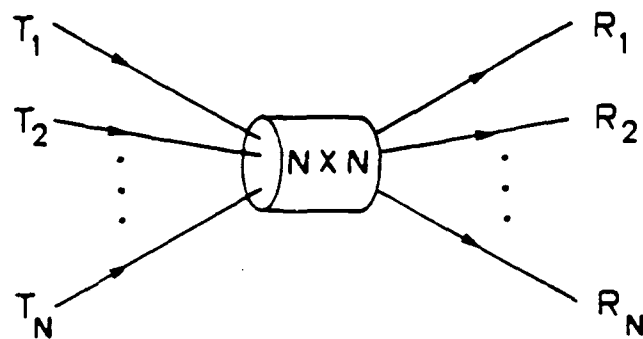
The third basic configuration we consider is the *tree* (T) configuration. This configuration differs from the star configuration by the fact that the data subnetwork is based on a *minimum-depth* binary-tree topology, as depicted in Fig. 4.4(c). (A linear topology can also be viewed as a binary-tree topology but with the maximum depth possible.) We require the collection and distribution subnetworks to be configured for minimum depth, in order to minimize the maximum number of couplers between transmitters and receivers.

4.2.2.2 Hybrid Configurations

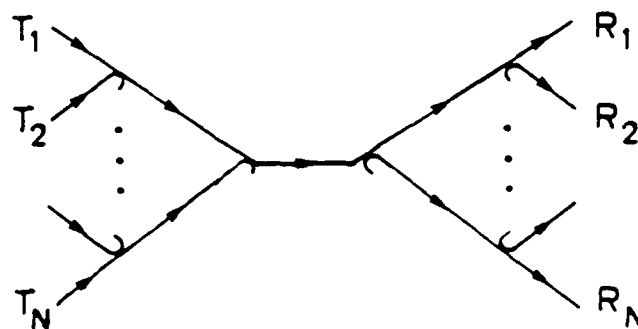
Among the hybrid configurations, we choose two which combine some of the features of the linear-bus, star, and tree configurations. In both of these configurations, the collection subnetwork has a linear-bus topology which also provides the control function. This choice is made to avoid the need for a separate control subnetwork. The distribution subnetwork of the first configuration is based on a star topology, as depicted in Fig. 4.5(a). (Note that only one of the star coupler's inputs is used.) This configuration is referred to as the *Linear-Star* (LS) configuration.



(a) Linear (L)

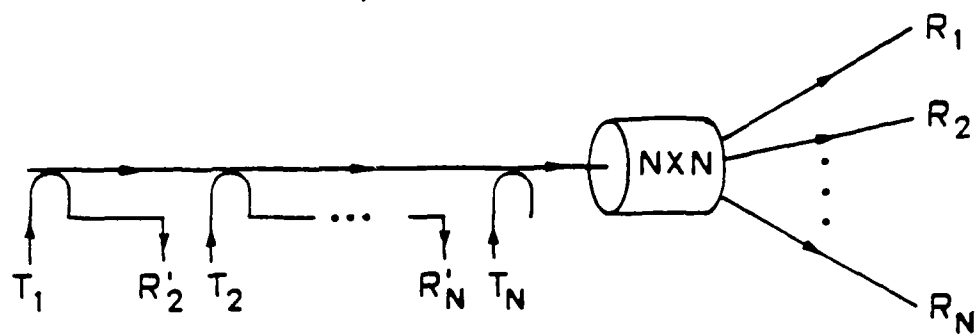


(b) Star (S)

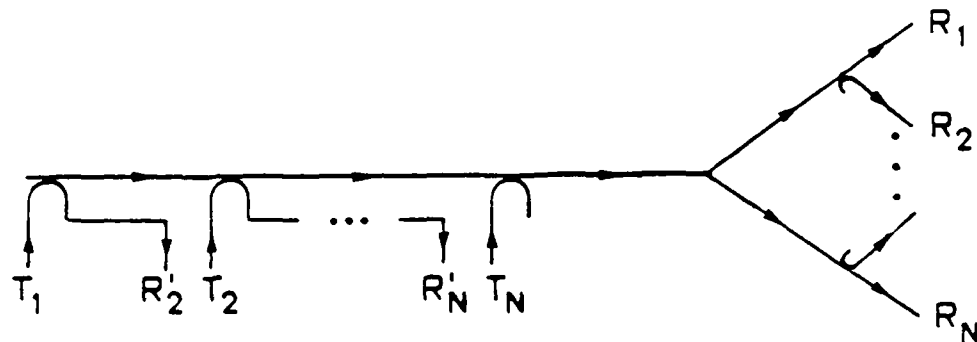


(c) Tree (T)

Fig. 4.4 Basic configurations for the data subnetwork.



(a) Linear-Star (LS)



(b) Linear-Tree (LT)

Fig. 4.5 Hybrid configurations for the data subnetwork.

The distribution subnetwork of the second configuration is based on a tree topology, as shown in Fig. 4.5(b). This configuration is referred to as the *Linear-Tree* (LT) configuration.

4.2.2.3 Compound Configurations

Among the many compound configurations one may conceive, we consider two in which the set of collection subnetworks and the set of distribution subnetworks all have the same topology. The first compound configuration is the *linear-star-linear* (LSL) in which the subnetworks have a linear-bus topology (see Fig. 4.6(a)). The second compound configuration is the *tree-star-tree* (TST) in which the subnetworks

have a minimum-depth binary-tree topology (see Fig. 4.6(b)). Both compound configurations require the provision of a linear control subnetwork.

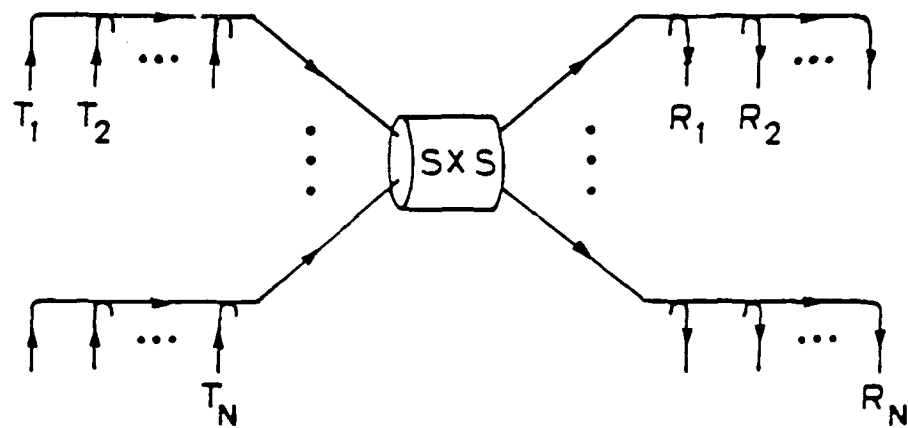
The basic configurations which are described above are special instances of the two compound configurations. More specifically, when $S = N$ both compound configurations are the same as the star configuration. The LSL configuration with $S = 1$ is nearly the same as the linear configuration. (The only difference is that the collection bus of the linear configuration also provides the control function, while the LSL configuration with $S = 1$ does not provide the control function. The two configurations, however, can support almost exactly the same number of stations. See Appendix A for details.) The TST configuration with $S = 1$ is the same as the tree configuration. We refer to S as the *funneling width*.

4.3 Power Budget Analysis

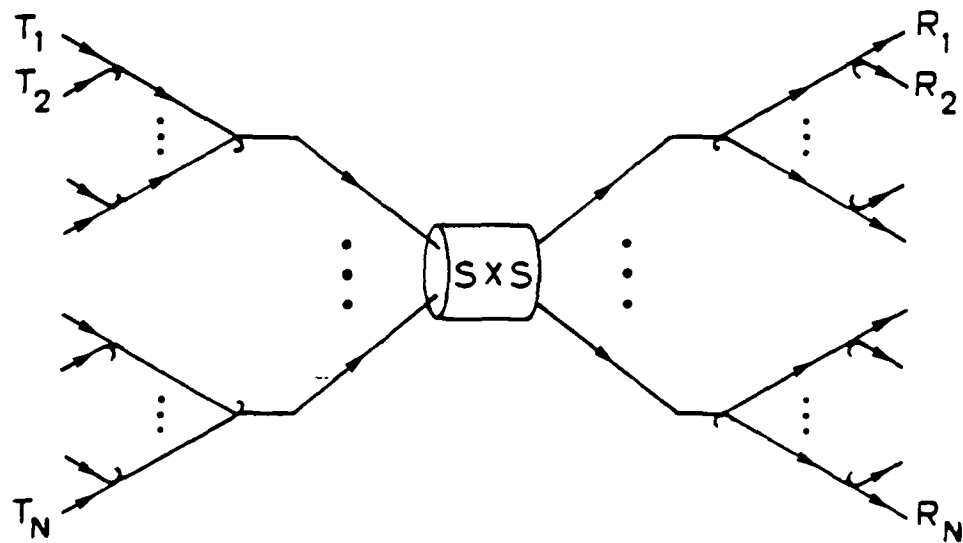
In this section, we provide a unified approach to the power budget analysis applicable to the various configurations considered in section 4.2. Using this approach, we can determine whether or not a configuration can support a given number of stations. In section 4.4, based on this approach, the maximum number of stations that each configuration can support is numerically computed. However, at the end of this section, we present a theoretical bound on the (maximum) number of stations in data subnetworks.

4.3.1 Formulation of the Problem

Let S denote the set of all transmitter and receiver pairs that need to communicate. Let M denote the number of couplers in the network. For $i = 1, 2, \dots, M$, let



(a) Linear-Star-Linear (LSL)



(b) Tree-Star-Tree (TST)

Fig. 4.6 Compound configurations for the data subnetwork.

C_i denote the i th coupler and x_i its coupling fraction. Let $\mathbf{x} \triangleq (x_1, x_2, \dots, x_M)$ and $\mathcal{X} \triangleq \{\mathbf{x} : x_{\min} \leq x_i \leq x_{\max}; i = 1, 2, \dots, M\}$. For a point A and a point B downstream of A , let $G(A, B; \mathbf{x})$ denote the transmittance from A to B as a function of $\mathbf{x}$. Note that $G(A, B; \mathbf{x})$ is simply the product of the transmittances of components through which the path from A to B passes.

For every transmitter T in the network, let $\mathcal{P}(T)$ denote its output power in watts. For every receiver (or sensor) R in the network, let $\mathcal{P}(R)$ denote its sensitivity in watts. To have full connectivity in the network, for some $\mathbf{x} \in \mathcal{X}$, we must have

$$\mathcal{P}(T)G(T, R; \mathbf{x}) \geq \mathcal{P}(R) \quad \forall (T, R) \in \mathcal{S} \quad (4.3)$$

Let $\mathcal{P}_T$ be the maximum $\mathcal{P}(T)$ among all transmitters. Let $\mathcal{P}_R$ be the minimum $\mathcal{P}(R)$ among all receivers. To simplify the analysis without loss of generality, we construct a virtual network as follows: we divide the transmittance of the fiber connecting a transmitter T to the network by a factor $\frac{\mathcal{P}_T}{\mathcal{P}(T)}$, and that of the fiber connecting a receiver R to the network by a factor $\frac{\mathcal{P}(R)}{\mathcal{P}_R}$. Now in terms of the transmittances in the virtual network, let

$$\underline{G}(\mathcal{S}; \mathbf{x}) \triangleq \min_{(T, R) \in \mathcal{S}} G(T, R; \mathbf{x}) \quad (4.4)$$

Then full connectivity is possible if

$$\max_{\mathbf{x} \in \mathcal{X}} \underline{G}(\mathcal{S}; \mathbf{x}) \geq \frac{\mathcal{P}_R}{\mathcal{P}_T} \quad (4.5)$$

4.3.2 Optimization of Couplers

Here, we present a technique for optimization of couplers which is applicable to any network configuration that satisfies the following two conditions:

C1: The configuration must contain optical paths *only* between those transmitter and receiver pairs which need communicate, namely the pairs in the set S .*

C2: The subnetworks connected to the ports of a coupler should be non-overlapping; this implies that the configuration is loop free.†

All configurations presented in the previous section above satisfy both conditions.

Our objective is to find $\mathbf{x}$ which maximizes $\underline{G}(\mathcal{S}; \mathbf{x})$. Consider some coupler C_i , $1 \leq i \leq M$, and fix the coupling fraction of every coupler C_j , other than C_i , to some value, say x_j . Next, consider the maximization of $\underline{G}(\mathcal{S}; \mathbf{x})$ over x_i . (This implies that we are performing a local optimization of $\underline{G}(\mathcal{S}; \mathbf{x})$ over x_i .) Clearly, x_i affects the transmittance from a transmitter T to a receiver R , only if the path from T to R goes through the coupler C_i . Let S_i denote the set of such (T, R) pairs. As a consequence of condition C1, a sufficient but not necessary condition for x_i to maximize $\underline{G}(\mathcal{S}; \mathbf{x})$ is that it maximizes $\underline{G}(S_i; \mathbf{x})$.

Based on condition C2, the coupler C_i partitions the configuration into four subnetworks, as depicted in Fig. 4.7. Let $\mathcal{N}_i^a$ denote the subnetwork connected to port C_i^a . If there is a (T, R) pair in $\mathcal{S}$ such that the path from T to R goes through C_i^a , let $G_i^a(\mathcal{S}; \mathbf{x})$ denote the minimum of $G(T, C_i^a; \mathbf{x})$ over all such pairs. Otherwise, let $G_i^a(\mathcal{S}; \mathbf{x})$ be infinity. Note that $G_i^a(\mathcal{S}; \mathbf{x})$ is not a function of x_i . (To simplify some of the following expressions, we may use G_i^a for $G_i^a(\mathcal{S}; \mathbf{x})$.) Similar entities are defined for the other ports of C_i , as specified in Fig. 4.7. Based on these definitions, and the coupler transmittance relations given by (4.1), we get

$$\underline{G}(S_i; \mathbf{x}) = \alpha \min\{G_i^a G_i^c (1 - x_i), G_i^a G_i^d x_i, G_i^b G_i^c x_i, G_i^b G_i^d (1 - x_i)\} \quad (4.6)$$

*For an example given below, this condition is somewhat relaxed.

†As demonstrated by example 2 below, the milder condition of just being loop free is not sufficient for the optimality of the conditions derived hereafter.

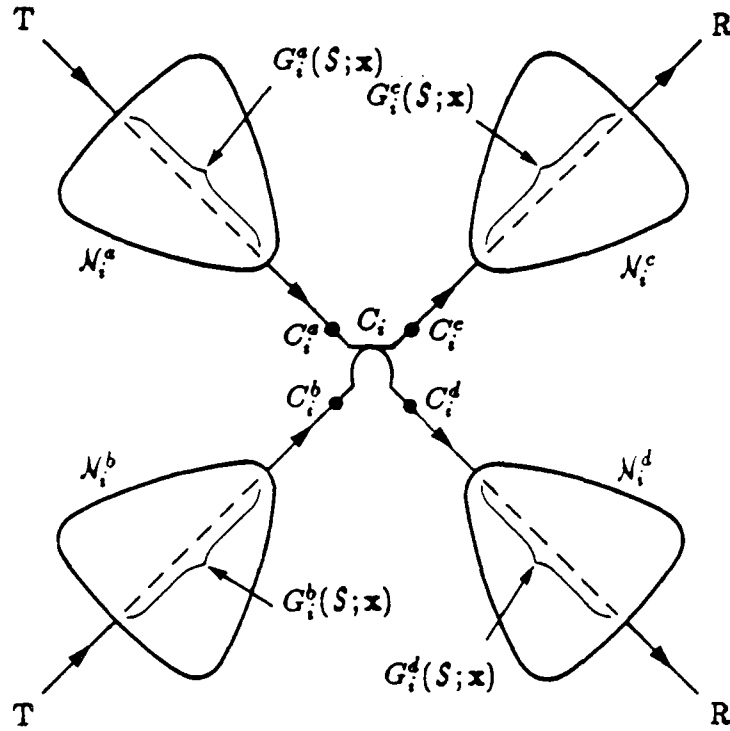


Fig. 4.7 The partitioning of a configuration by a coupler C_i and the corresponding notation.

A plot of $\underline{G}(S_i; \mathbf{x})$ as a function of x_i is given in Fig. 4.8. Let

$$[x]_{x_{\min}}^{x_{\max}} \triangleq \max(\min(x, x_{\max}), x_{\min}). \quad (4.7)$$

From Fig. 4.8, we observe that $\underline{G}(S_i; \mathbf{x})$ is maximized at

$$x_i = \left[\frac{\min\{G_i^a G_i^c, G_i^b G_i^d\}}{\min\{G_i^a G_i^c, G_i^b G_i^d\} + \min\{G_i^a G_i^d, G_i^b G_i^c\}} \right]_{x_{\min}}^{x_{\max}} \quad (4.8)$$

Up to now, we have been assuming that all the coupling fractions other than x_i are fixed. Now, we select all the coupling fractions in such a way that for every $i = 1, 2, \dots, M$ the local optimality condition given by (4.8) is satisfied.

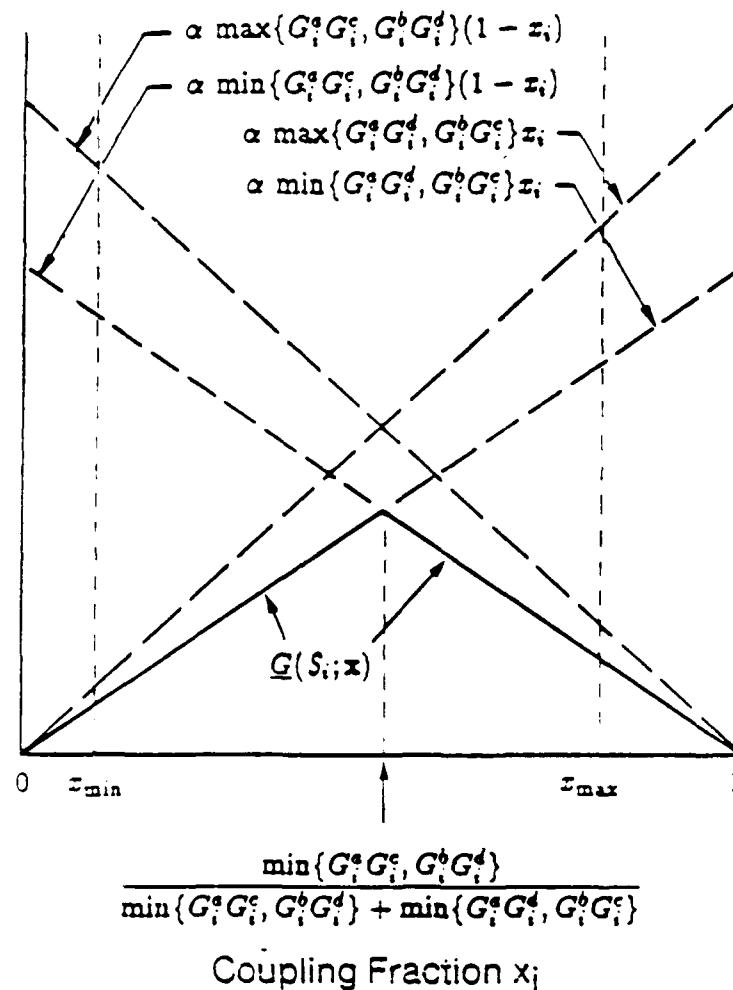


Fig. 4.8 $\underline{G}(S; \mathbf{x})$ as a function of x_i .

As proved in Appendix A, this set of coupling fractions also results in the maximum $\underline{G}(S; \mathbf{x})$.* Now the technique for maximizing $\underline{G}(S; \mathbf{x})$ is clear: to find an optimum set of coupling fractions, we solve the set of equations given by (4.8) for $i = 1, 2, \dots, M$.

Claim: If $\mathbf{x}$ satisfies (4.8) for every $1 \leq i \leq M$, then it maximizes $\underline{G}(S; \mathbf{x})$.

Proof: Let $\mathbf{x}$ be a vector satisfying (4.8). The proof consists of showing that if

*Example 1 below shows that (4.8) is not a necessary condition for optimality.

there exists a y in $\mathcal{X}$ such that $\underline{G}(S; y) > \underline{G}(S; x)$, then there must be a infinite sequence of sets S^n , $n = 0, 1, 2, \dots$ such that $S^n \neq \emptyset$, $S^0 = S$, and $S^{n+1} \subset S^n$. But since S is finite, this is a contradiction.

Let $y^0 = y$. By assumption $\underline{G}(S^0; y^0) > \underline{G}(S^0; x)$ and x satisfies (4.8). Now, we show that if there exists a pair (S^n, y^n) such that

$$A_n: \quad \underline{G}(S^n; y^n) > \underline{G}(S^n; x), \text{ and}$$

$$B_n: \quad x \text{ satisfies (4.8) for every coupler } C_i \text{ for which } x_i \neq y_i^n,$$

then there exists a pair (S^{n+1}, y^{n+1}) satisfying A_{n+1} and B_{n+1} , such that

$$C_{n+1}: \quad S^{n+1} \neq \emptyset \text{ and } S^{n+1} \subset S^n.$$

There must be a (T_n, R_n) in S^n such that $G(T_n, R_n; x) = \underline{G}(S^n; x)$. From A_n , we have $G(T_n, R_n; y^n) > G(T_n, R_n; x)$. Therefore, there must be a coupler C_{i_1} on the path from T_n to R_n such that its transmittance according to y^n is strictly greater than its transmittance according to x . Without loss of generality, assume these transmittances are $y_{i_1}^n$ and $x_{i_1}^n$, respectively (Another possibility is $1 - y_{i_1}^n$ and $1 - x_{i_1}^n$, respectively); we have $y_{i_1}^n > x_{i_1}^n$.

Since $y_{i_1}^n \leq x_{\max}$ we have $x_{i_1} < x_{\max}$. Based on B_n , there exists a (T_{n+1}, R_{n+1}) in S^n such that the transmittance from T_{n+1} to R_{n+1} according to x is $1 - x_{i_1}$ and $G(T_{n+1}, R_{n+1}; x) = G(T_n, R_n; x) = \underline{G}(S^n; x)$ (See Fig. 4.8.) Without loss of generality, we shall assume that T_{n+1} is connected to $C_{i_1}^a$ and R_{n+1} is connected to $C_{i_1}^c$. (Another possibility is T_{n+1} connected to $C_{i_1}^b$ and R_{n+1} connected to $C_{i_1}^d$.)

Let S^{n+1} be the set of all (T, R) in S^n , such that T is in either $\mathcal{N}_{i_1}^a$ or $\mathcal{N}_{i_1}^c$ subnetwork and R is in either $\mathcal{N}_{i_1}^a$ or $\mathcal{N}_{i_1}^c$ subnetwork. Since $(T_{n+1}, R_{n+1}) \in S^{n+1}$ and $(T_{n+1}, R_{n+1}) \notin S^n$, C_{n+1} follows.

Let y^{n+1} be such that $y_i^{n+1} = y_i^n$ if C_i is in $\mathcal{N}_{i_1}^a$ or $\mathcal{N}_{i_1}^c$, and $y_i^{n+1} = x_i$ otherwise. Since $1 - y_{i_1}^{n+1} > 1 - y_{i_1}^n$, $\underline{G}(S^n; y^{n+1}) > \underline{G}(S^n; y^n)$. Therefore A_{n+1} follows.

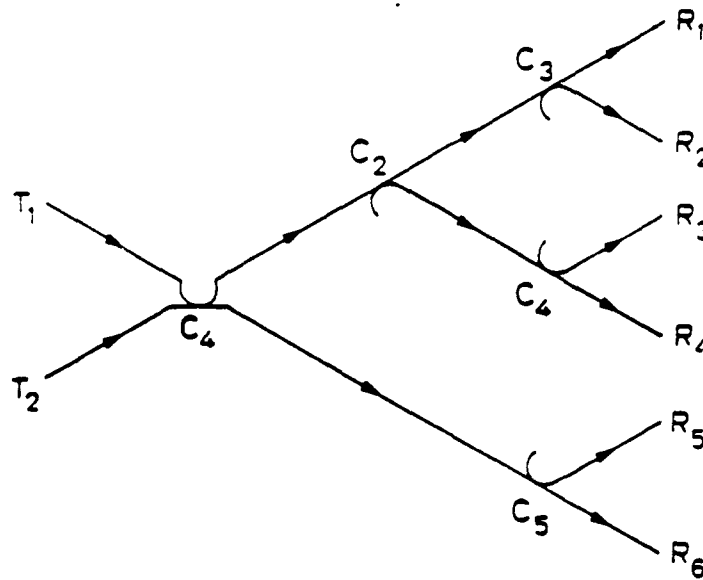


Fig. 4.9 A configuration in which local optimality of the couplers is not necessary for global optimality.

Note that, for every coupler C_i in N_i^a or N_i^c , $G_i^w(S^{n-1}; x) = G_i^w(S^n; x)$, for $w = a, b, c, d$. Therefore from B_n , B_{n+1} follows. Q.E.D

Example 1. This example shows that (4.8) is not a necessary condition for optimality: consider the configuration shown in Fig. 4.9. We assume that all components are ideal. Then $x = (\frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \frac{1}{2})$ satisfies (4.8) and $\underline{G}(S; x) = \frac{1}{8}$. However, $\underline{G}(S; y) = \frac{1}{8}$ for any $y = (\frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \frac{1}{2}, y_5)$ where $\frac{1}{4} \leq y_5 \leq \frac{3}{4}$.

Example 2. This example shows that (4.8) is not a sufficient condition for optimality in a configuration which is loop free but does not satisfy C1: consider the configuration shown in Fig. 4.10. We assume that all components are ideal. (4.8) is satisfied for any $x = (x, x, x, x)$ where $0 \leq x \leq 1$. However, $\underline{G}(S; x)$ is only maximized at $x = (\frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \frac{1}{2})$.

For a general configuration, the set of equations given by (4.8), for $i = 1, 2, \dots, M$, is solved iteratively. However, in the case of a collection or distribution subnetwork

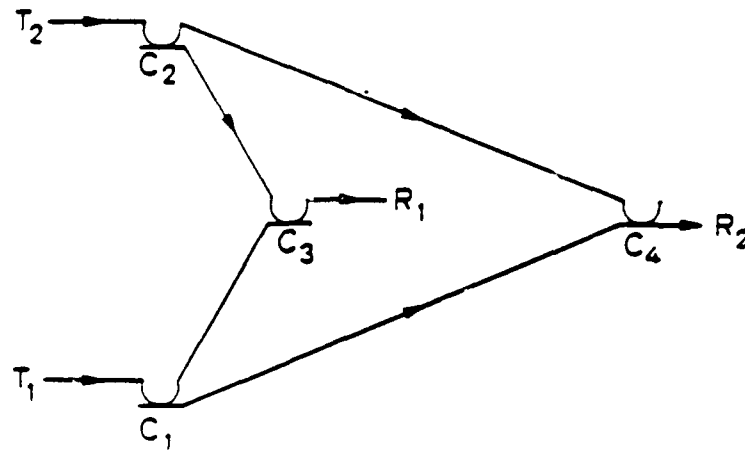


Fig. 4.10 A loop-free configuration which does not satisfy condition C2.

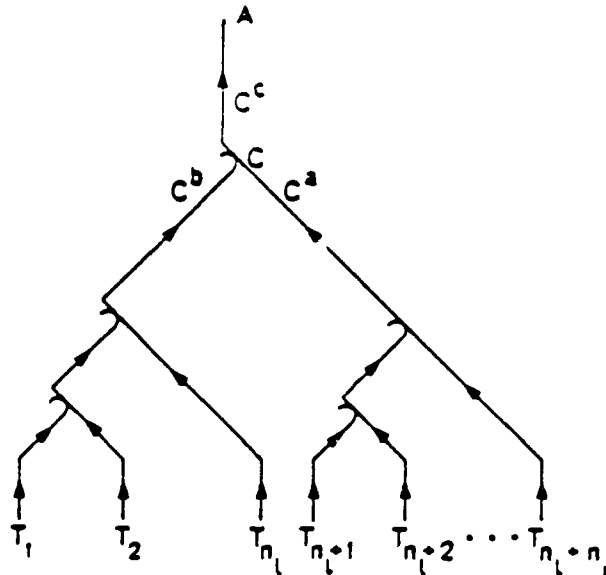


Fig. 4.11 A collection subnetwork based on a (general) binary tree topology.

based on a binary tree topology, of which the linear-bus and the minimum-depth binary tree are special cases, we show that the solution can be obtained recursively. Consider a collection subnetwork based on the binary tree topology as depicted in Fig. 4.11. (Although, only the case of a collection subnetwork is discussed, similar

statements hold for the case of a tree distribution subnetwork. From the above definitions, for every coupler C_i in this configuration, G_i^d is infinity. Therefore, (4.8) reduces to

$$x_i = \left[\frac{G_i^a}{G_i^a + G_i^b} \right]_{x_{\min}}^{x_{\max}} \quad i = 1, 2, \dots, M \quad (4.9)$$

Using (4.9), one can optimize the couplers in such an order that, for every $i = 1, 2, \dots, M$, when C_i is optimized, both $G_i^a(S; \mathbf{x})$ and $G_i^b(S; \mathbf{x})$ can be computed in terms of the coupling fractions which are already optimized. First, one optimizes a coupler which has each of its input ports connected to a transmitter; in each subsequent step, one optimizes a coupler which has each of its input ports connected to either a transmitter or to an already optimized coupler. It is interesting to note that, when $x_{\min} = 0$ and $x_{\max} = 1$, the optimization results in the same transmittance from every transmitter to the funneling point. Furthermore, if all the other components are ideal - i.e., the excess-loss of couplers, the insertion loss of connectors and splices, and the fiber attenuation are all zero - then this transmittance is exactly equal to $\frac{1}{n}$, where n is the number of transmitters connected to the collection subnetwork. We show this by induction: clearly the assertion is true for $n = 1$. For $n \geq 2$, assume it is true for any n' less than n . In Fig. 4.11, let n_l and n_r , ($n_l + n_r = n$) denote the number of stations in the left and right subtrees, respectively. Let x be the coupling fraction of coupler C which is at the root of the tree. By assumption $G(T_i, C^b) = \frac{1}{n_l}$ for $1 \leq i \leq n_l$ and $G(T_i, C^a) = \frac{1}{n_r}$ for $n_l + 1 \leq i \leq n_r$. From (4.9) we have

$$x = \frac{\frac{1}{n_r}}{\frac{1}{n_l} + \frac{1}{n_r}} = \frac{n_l}{n} \quad (4.10)$$

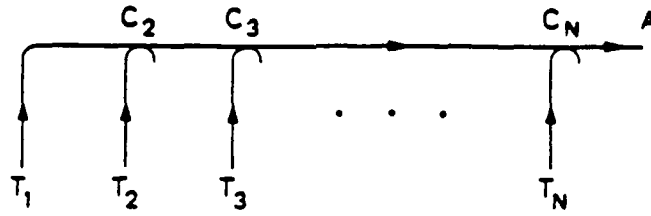


Fig. 4.12 Linear collection subnetwork.

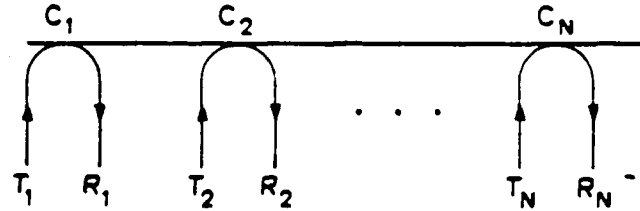


Fig. 4.13 A configuration which does not satisfy condition C1.

and

$$\begin{aligned} G(T_i, C^c) = xG(T_i, C^b) &= \frac{1}{n} & 1 \leq i \leq n_l \\ G(T_i, C^a) = (1-x)G(T_i, C^a) &= \frac{1}{n} & n_l + 1 \leq i \leq n \end{aligned} \quad (4.11)$$

In a linear collection subnetwork with ideal components where the couplers are indexed as in Fig. 4.12, the optimum coupling fractions are given by

$$x_i = \frac{1}{i} \quad 2 \leq i \leq N \quad (4.12)$$

Finally, we consider the configuration shown in Fig. 4.13, which has been extensively treated in the literature [ALTM77, AURA77, VILLS1]. This configuration does not satisfy condition C1: although the transmitter and receiver of each station need not communicate, there exists a path between them. Fortunately, (4.8) can

easily be modified for this configuration as follows. From Fig. 4.7, it is clear that, for $i = 1, 2, \dots, N$, T_i is the only transmitter in network N_i^b and R_i is the only receiver in network N_i^d . Hence, to disregard the path from T_i to R_i in the optimization, we simply drop the last of the four terms which appear between the brackets in (4.6). Then, the new set of conditions for local optimality becomes

$$x_i = \left[\frac{G_i^a G_i^c}{G_i^a G_i^c + \min\{G_i^a G_i^d, G_i^b G_i^c\}} \right]_{x_{\min}}^{x_{\max}} \quad i = 1, 2, \dots, N \quad (4.13)$$

Based on a proof similar to that given in the appendix A, we can show that when the coupling fractions satisfy (4.9), they also maximize $\underline{G}(S; \mathbf{x})$. It is interesting to note that for the ideal components, when the couplers are indexed as in Fig. 4.13, the optimum coupling fractions are

$$x_i = \max\left(\frac{1}{i}, \frac{1}{N-i+1}\right) \quad (4.14)$$

The optimum coupling fraction is equal to 1 for the end stations and decreases as we move toward the middle of the bus. For non-ideal components, the optimum coupling fractions exhibit the same behavior [ALTM77, AURA77]. (See also Appendix A.)

4.3.3 A Bound on Number of Stations

Here, we derive a theoretical bound on the (maximum) number of stations in a data subnetwork. To derive this bound we assume the best possible situation, i.e. ideal components and individually optimized couplers. Furthermore, we assume that all transmitters have the same output power $\mathcal{P}_T$, and all receivers have the same sensitivity $\mathcal{P}_R$, in Watts. Consider a compound data subnetwork as shown in Fig. 4.14 in which the collection and distribution subnetworks are based on a

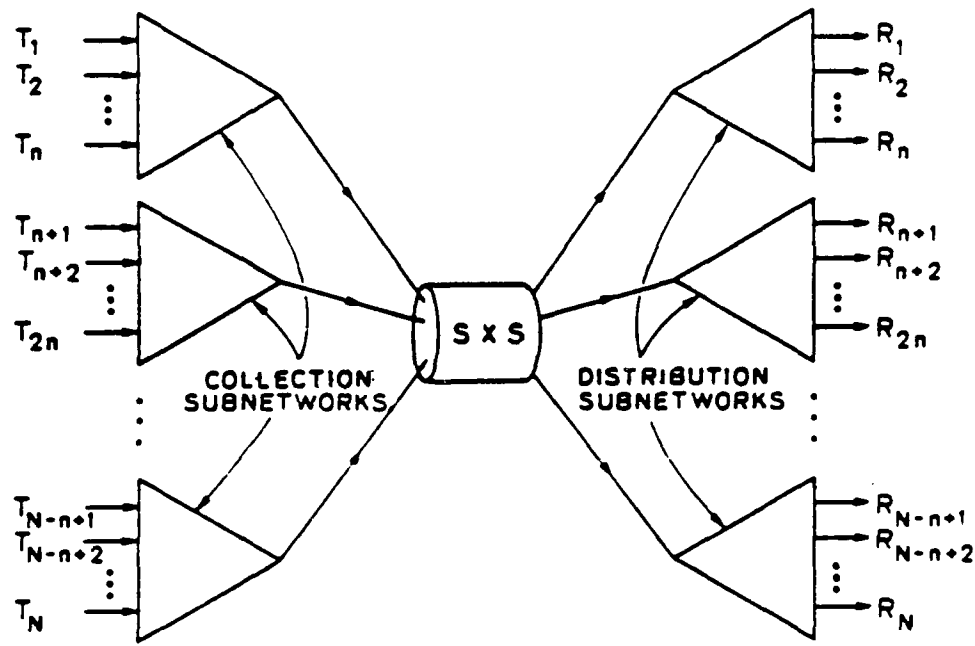


Fig. 4.14 A general compound data subnetwork.

(general) binary-tree topology. To keep the bound simple, we assume that the number of stations N is a multiple of S ; i.e., $N = nS$. Note that the LSL and TST data subnetworks – and hence the L, T, and S data subnetworks – are special cases of this data subnetwork.

From (4.11), we know that the transmittance from every transmitter to the star input to which the transmitter is connected, is $\frac{1}{n}$. Similarly, the transmittance from every star output to every receiver to which the output is connected, is $\frac{1}{n}$. Furthermore, the transmittance across an ideal star is $\frac{1}{S}$. Therefore, the transmittance from every transmitter to every receiver is $\frac{1}{Sn^2}$ or $\frac{S}{N^2}$. In order to have full connectivity in the data subnetwork the product of this transmittance and the *power margin* $\mathcal{M} \triangleq \frac{P_T}{P_R}$ must be greater than one, or equivalently,

$$N \leq \sqrt{S\mathcal{M}} \quad (4.15)$$

The equality can only be achieved when the components are ideal and the couplers

are individually optimized (and $\sqrt{SM}$ is an integer.) In L and T subnetworks, S is equal to 1. Therefore,

$$N \leq \sqrt{M} \quad (4.16)$$

It is easy to show that (4.16) also holds for LT and LS data subnetworks. In the S subnetwork, the funneling width S is equal to N . Therefore

$$N \leq M \quad (4.17)$$

Based on (4.15), for ideal components and individually optimized couplers, the maximum number of stations $N_{\max}$ as function of S is plotted in Fig. 4.15. The left-most and the right-most points on every constant-power-margin line-segment correspond to the L and S data subnetworks, respectively. Clearly, as S increases more stations can be supported. However, this is achieved at the cost of longer fiber. Therefore it is desirable to know how one can trade-off fiber length with maximum number of stations.

We assume that the stations are equally spaced along a straight line. Furthermore, we normalize the total fiber length to twice the distance between the first and the last station. Hence, the normalized fiber length for the star configuration is exactly N . The normalized fiber length in the LSL configuration is given by

$$l = \frac{(S+1)N - 2S}{2(N-1)} + \frac{S}{2} \quad (4.18)$$

In the case of the TST configuration, assuming that N is a power of 2, the normalized fiber length is given by

$$l = \frac{N \left[S - 1 + \log_2 \frac{N}{S} \right]}{2(N-1)} + \frac{S}{2} \quad (4.19)$$

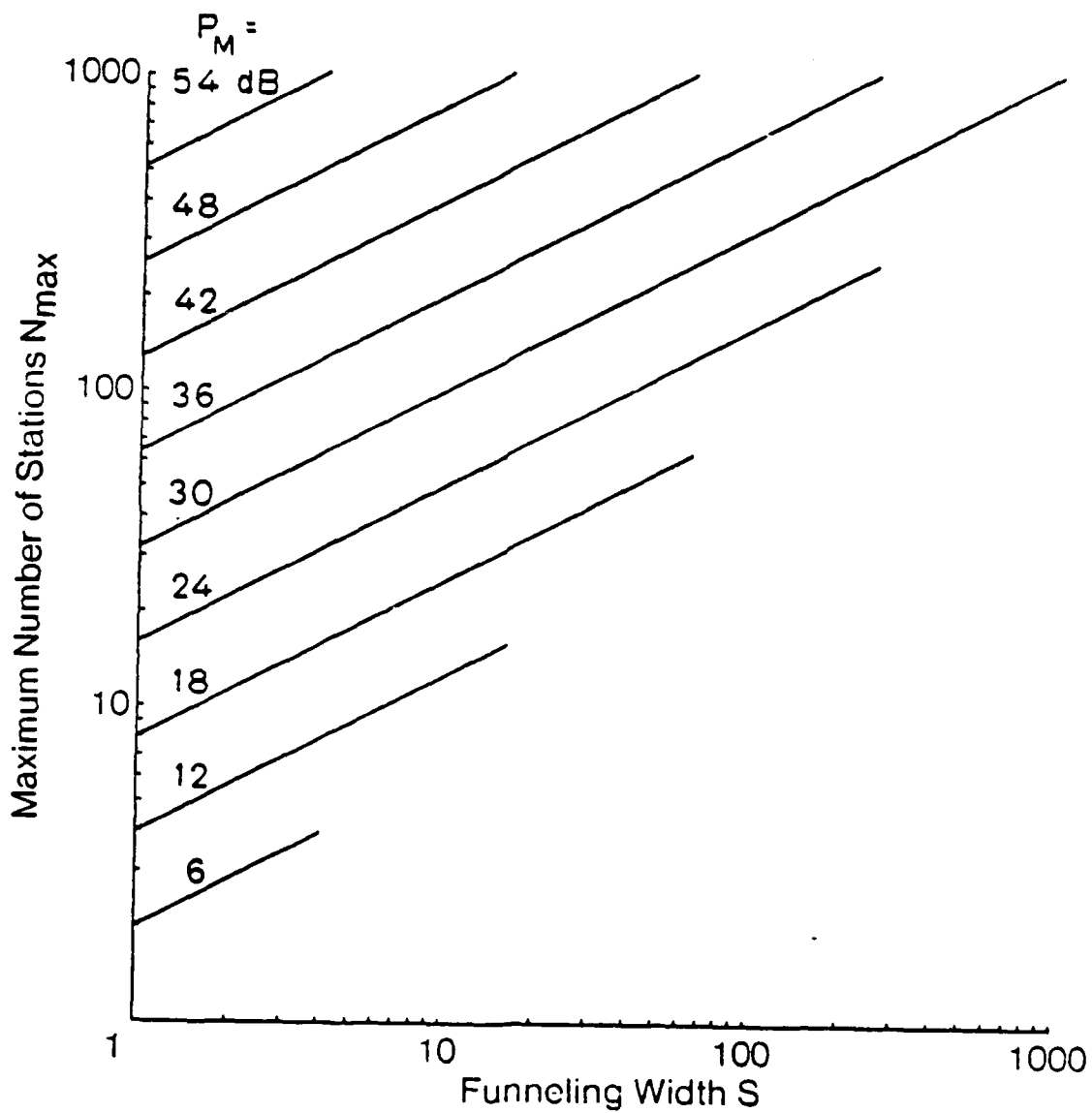


Fig. 4.15 The maximum number of stations as a function of the funneling width in a general data subnetwork. It is assumed that the components are ideal and the couplers are individually optimized.

The trade-off between maximum number of stations and fiber length is shown in Fig. 4.16. The solid and dashed lines correspond to the TST and LSL data subnetworks, respectively.

4.3.4 Coupler Optimization in linear subnetworks

In this section, we show how the coupling coefficients in the linear subnetworks are optimized. Here, these subnetworks are divided into three groups: linear collection and distribution data subnetworks, stretched control and collection-and-control subnetworks, and bypassed control subnetwork. The normalized fiber lengths in the compound configurations are also given.

A. Optimum Couplers in Linear Collection and Distribution Data Subnetworks: First, consider the distribution subnetwork of the linear configuration shown in Fig. 4.4(a). Let x_i denote the coupling coefficient of the i th receive coupler. For $1 \leq i \leq N-1$, define H_i as the minimum transmittance from the bus output of the i th coupler to all downstream receivers. Define H_N to be infinity. For symmetric couplers we have

$$H_{i-1} = \alpha \min\{(1 - x_i)H_i, x_i\} \quad 2 \leq i \leq N \quad (4.20)$$

H_{i-1} is maximized at

$$x_i = \frac{H_i}{1 + H_i} \quad 1 \leq i \leq N \quad (4.21)$$

Hence, the optimum coefficients can be recursively computed using (4.20) and (4.21). For asymmetric couplers we have

$$H_{i-1} = \alpha \min\{x_i H_i, 1 - x_i\} \quad 2 \leq i \leq N \quad (4.22)$$

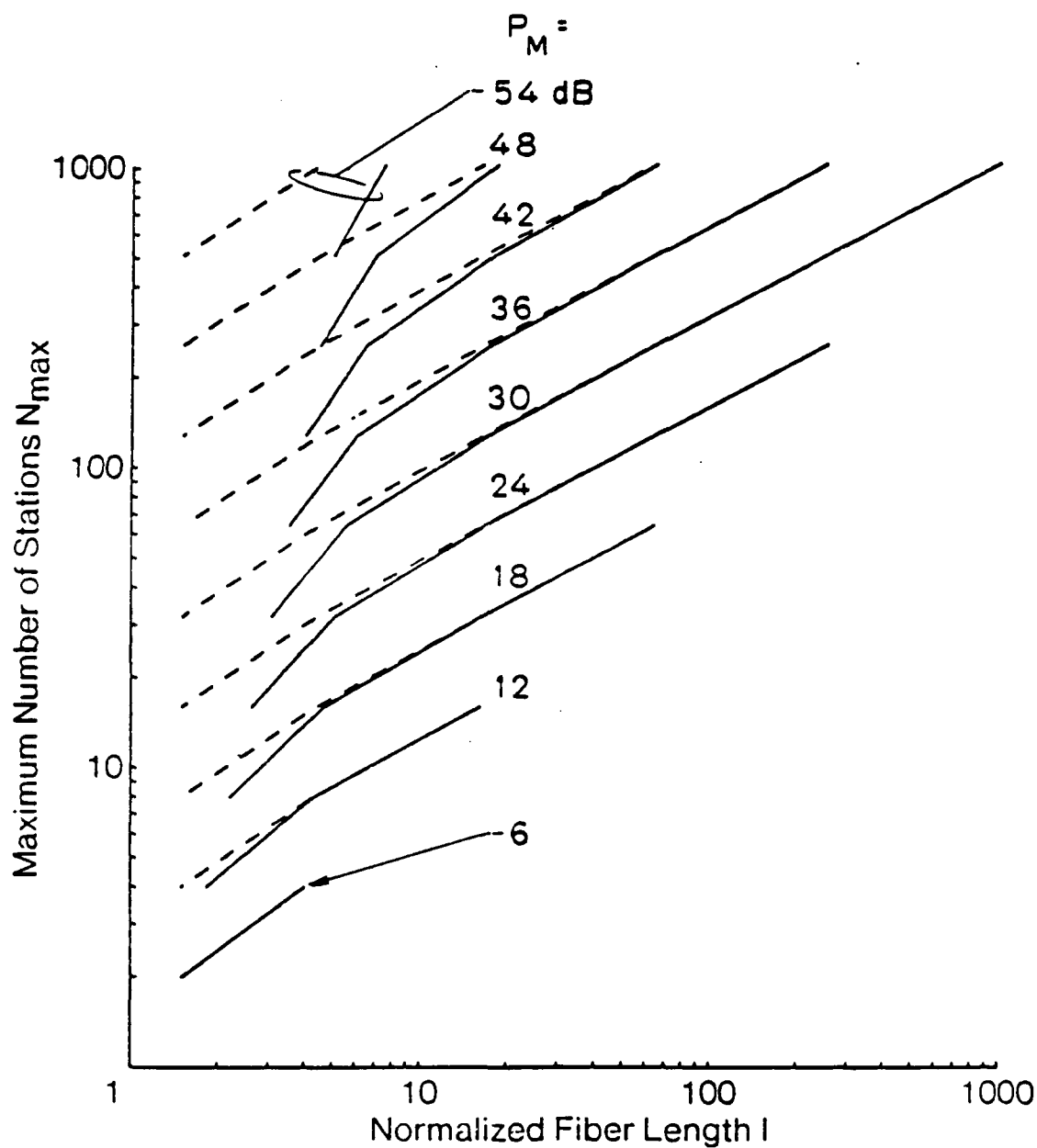


Fig. 4.16 The trade-off between the maximum number of stations and fiber length in TST (solid lines) and LSL (dashed lines) data subnetworks. It is assumed that the components are ideal and the couplers are individually optimized.

and

$$x_i = \frac{1}{1 + H_i} \quad 1 \leq i \leq N \quad (4.23)$$

In the case of collection subnetworks (e.g. those in LSL as shown in Fig. 4.6(a)) and symmetric couplers, the coefficients are optimized as in the distribution subnetwork. For asymmetric couplers, the optimum coefficients are equal to unity.

B. Optimum Couplers in Stretched Control and Collection-and-Control Subnetworks: Consider the control subnetwork of Fig. 4.2(b) and the collection-and-control subnetwork of Fig. 4.4(a). Let H_i denote the minimum transmittance from the bus output of the i th transmit coupler to all downstream receivers. Similarly, let F_i denote the minimum transmittance from all upstream transmitters to the bus input of the i th transmit coupler. H_N in the collection-and-control subnetwork is equal to the minimum transmittance from the funneling point to receivers, and in the control subnetwork, it is equal to infinity. F_1 in both cases is equal to infinity. For symmetric couplers H_i is given by (4.20) and F_i is given by

$$F_{i+1} = \alpha \min\{(1 - x_i)F_i, x_i\} \quad 1 \leq i \leq N - 1 \quad (4.24)$$

Let x_i denote the coefficient of the i th transmit coupler. Based on (4.8), for symmetric couplers, the optimum coefficient are given by

$$x_i = \frac{\min\{1, F_i H_i\}}{\min\{1, F_i H_i\} + \min\{F_i, H_i\}} \quad 1 \leq i \leq N \quad (4.25)$$

For asymmetric couplers, x_N is equal to 1, and the other coefficients are given by (4.22) and (4.23).

C. Optimum Couplers in Bypassed Control Subnetwork: Consider the bypassed control subnetwork shown in Fig. 4.2(c). The required power margin is minimized

by optimizing the couplers on both the main and bypass fibers. But due to the presence of multiple signal paths between most transmitter-receiver pairs, joint optimization of main-fiber and bypass-fiber coupling coefficients is computationally complex. Instead we compute a set of suboptimum coefficients as follows.

For any coupler i on the main fiber, we let F_i and H_i be defined as in part B above. Consider first symmetric couplers. We optimize the coupling coefficients on a segment j of the main fiber in isolation using (4.20), (4.24), and (4.25) in terms of boundary values $F_{(j-1)n+1}$ and $H_{j,n}$. To avoid the computational complexity involved in a global optimization of the entire network, we make the approximation that $F_{(j-1)n+1}$ be equal to $H_{j,n}$, for all j . Furthermore, we numerically optimize $H_{j,n}$ in order to maximize the minimum transmittance from A_j to the segment's receivers and from the segment's transmitters to B_j (see Fig. 4.2(c)). For asymmetric couplers, since the transmittance from A_j to $R'_{j,n+1}$ is equal to $\alpha^n(1 - x_{j,n}) \prod_{i=(j-1)n+1}^{j'n-1} x_i$ and from $T_{(j-1)n+1}$ to B_j is equal to $\alpha^n \prod_{i=(j-1)n+1}^{j'n} x_i$, assuming again $F_{(j-1)n+1} = H_{j,n}$, the optimum $x_{j,n}$ is $\frac{1}{2}$. Equations (4.22) and (4.23) are used to compute x_i , for $(j-1)n+1 \leq i \leq jn-1$.

When the main-fiber coupling coefficients are optimized as above for every segment j , the minimum transmittance from A_j to the segment's receivers is equal to that from the segment's transmitters to B_j . Furthermore since the stations are assumed to be evenly divided among the segments, this minimum transmittance is the same for all segments. Therefore, the bypass coefficients can be optimized as in the stretched control configurations of Fig. 4.2(b), where T_j is replaced by B_j and R'_j is replaced by A_j .

4.4 Performance

4.4.1 Configurations for Control Subnetwork

In deriving the performance results, we assume that the factor limiting the maximum number of stations is the power margin and not the dynamic range, since dynamic ranges as large as 60 dB may be obtained with automatic-gain control (AGC) by incurring a small penalty in the power margin [MUOI84]. (Of course, a sufficiently long packet preamble for AGC and bit timing recovery is always required.) We ignore the signal attenuation due to cable loss: in the local environment, distances will usually be less than 1 kilometer between active repeaters and thus the fiber attenuation will be less than a few dB. For simplicity, the splice and/or connector insertion losses associated with the couplers are included in their excess loss.*

The coupling coefficients are individually optimized, as outlined in section 4.3.† Discretizing the coupling coefficients to a finite set of values reduces the maximum number of stations for a given power margin. For example in the basic control configuration, when the coupling coefficients are quantized to multiples of -4dB, the maximum number of stations is decreased by about 20% [LIMB84].

Figure 4.17 depicts the maximum number of stations as a function of the power margin P_M in the control subnetworks for a coupler excess loss of 1 dB. As it can be seen, the stretch feature nearly doubles the number of stations that can be supported for a given power margin. As mentioned above, for sufficiently large N , the bypass feature further reduces the required power margin. The use of asymmetric couplers results in an improvement in the performance of all the configurations and it is most significant in the case of the bypassed control.

*A more precise formulation of this problem, which takes into account fiber attenuation and splice and connector insertion losses, can be found in Appendix A.

†When all couplers are required to have the same coefficient, the maximum number of stations is significantly smaller than that when they are individually optimized (see Appendix A).

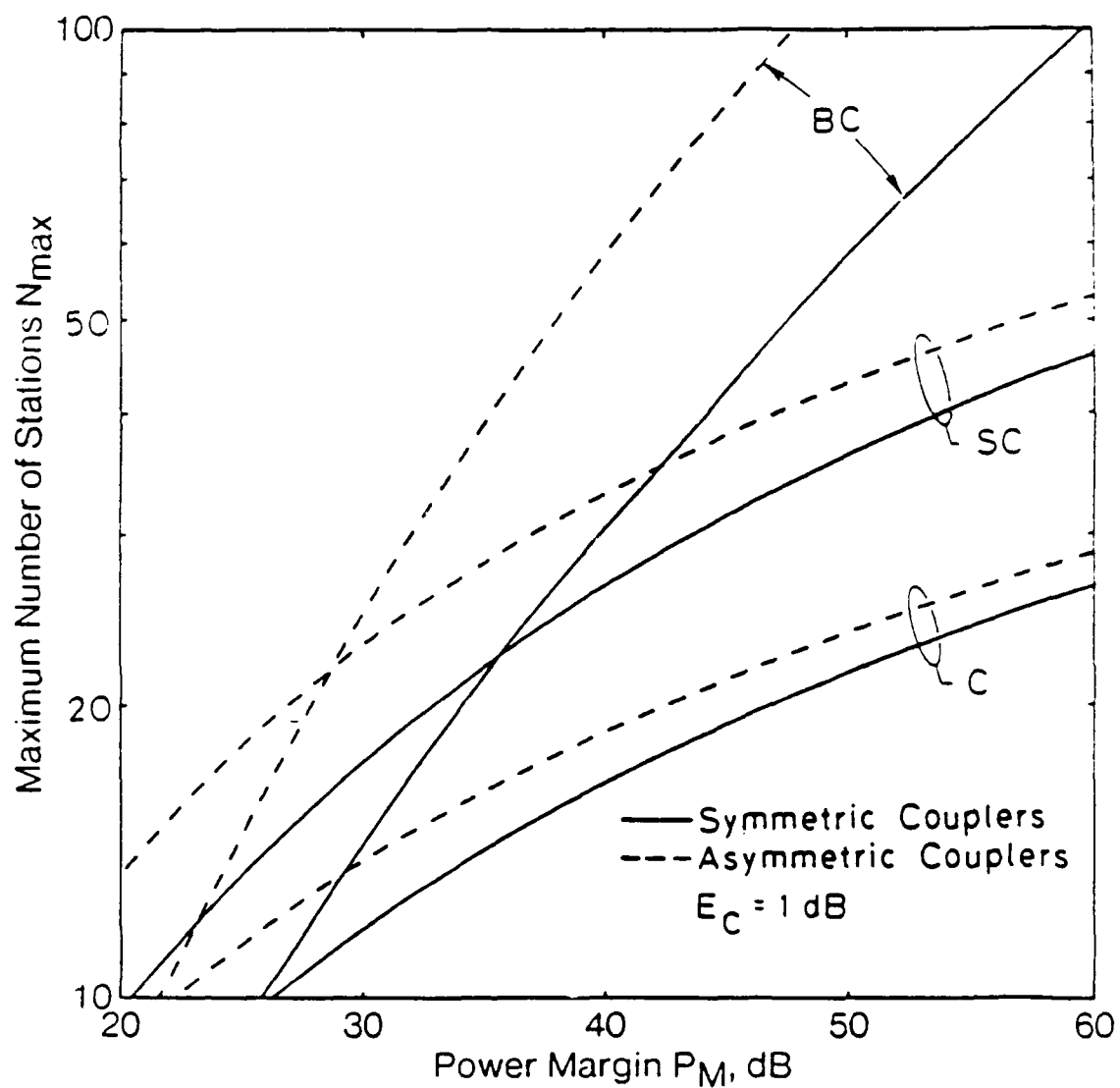


Fig. 4.17 Maximum number of stations versus power margin in the control sub-network for a coupler excess loss of 1dB.

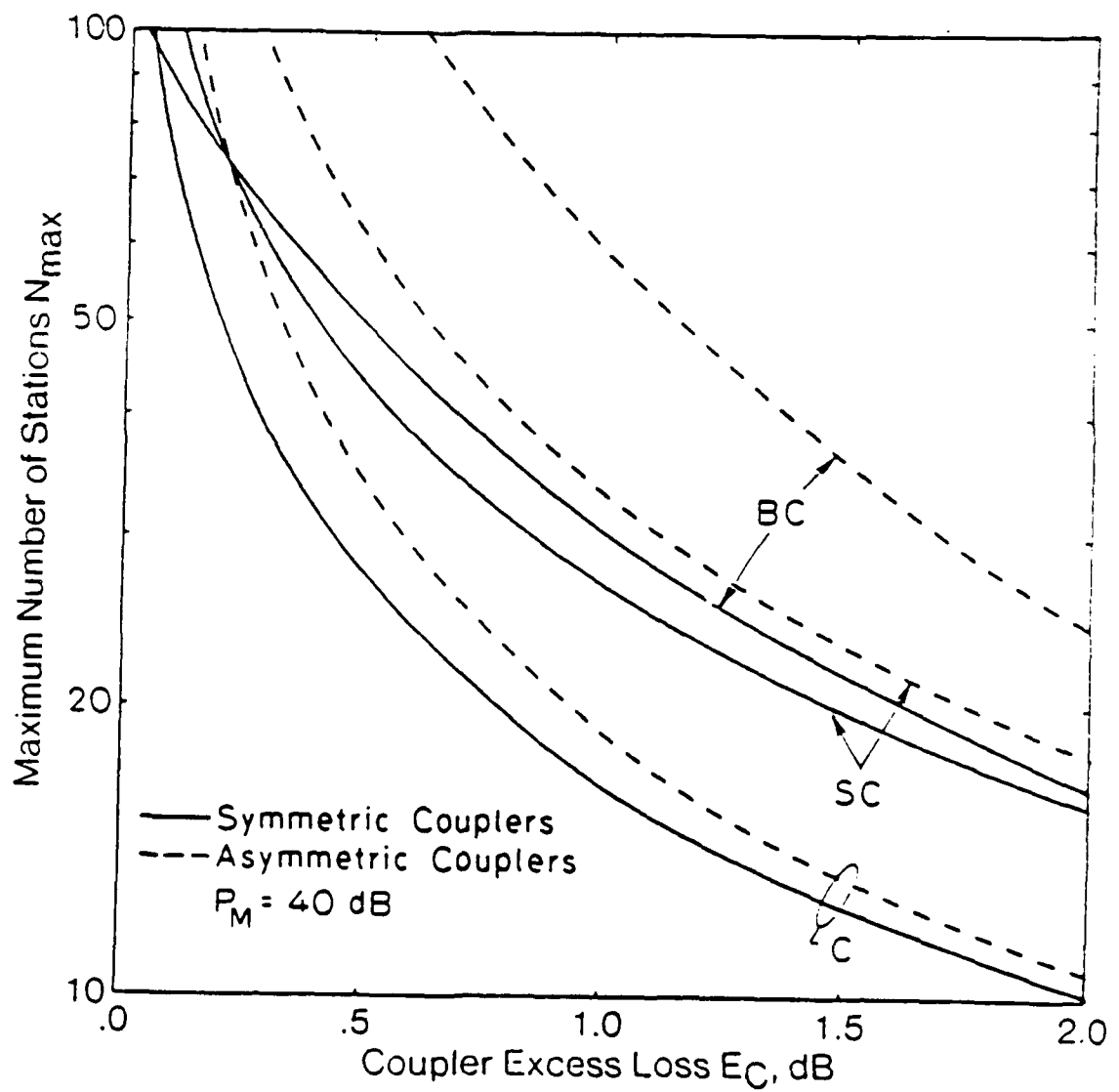


Fig. 4.18 Maximum number of stations versus coupler excess loss in the control subnetwork for a power margin of 40 dB.

Figure 4.18 shows the maximum number of stations as a function of the excess loss per coupler E_C , for a power margin of 40 dB. Although the maximum number of stations drops sharply as the coupler excess loss increases, the relative performance of all control configurations remains the same for nearly the entire range of coupler excess loss.

4.4.2 Configurations for Data Subnetwork

The maximum number of stations as a function of the power margin for the basic and hybrid data configurations is shown in Fig. 4.19. The excess loss of the 2×2 couplers, E_C , is assumed to be 1 dB. The excess loss of an $S \times S$ star coupler is assumed to be $E_C \log_2 S$, which is a good approximation for a family of commercially available star couplers [CANSTA], and is exact when a modular star based on 2×2 couplers [MARH84] is used (for S a power of two). In the linear subnetworks, the coupling coefficients are individually optimized as outlined in 4.3.4. Although the coupling coefficients for the tree subnetworks are also individually optimized, using $\frac{1}{2}$ for all the coupling coefficients results in an additional loss of only a few dB (see Appendix A). (In fact $\frac{1}{2}$ is optimum for all coupling coefficients when the number of stations is a power of two.)

As shown in Fig. 4.19, there exist considerable performance differences among the data configurations considered. (Not shown in this figure is a slight difference between the performance of the LS and LT configurations when the number of stations is not a power of two. See Appendix A.) These performance differences are caused by the variations in the reciprocity and excess losses incurred in the configurations.

First, we note that it is clearly the reciprocity loss that accounts for the performance disadvantage of the linear and hybrid configurations based on symmetric

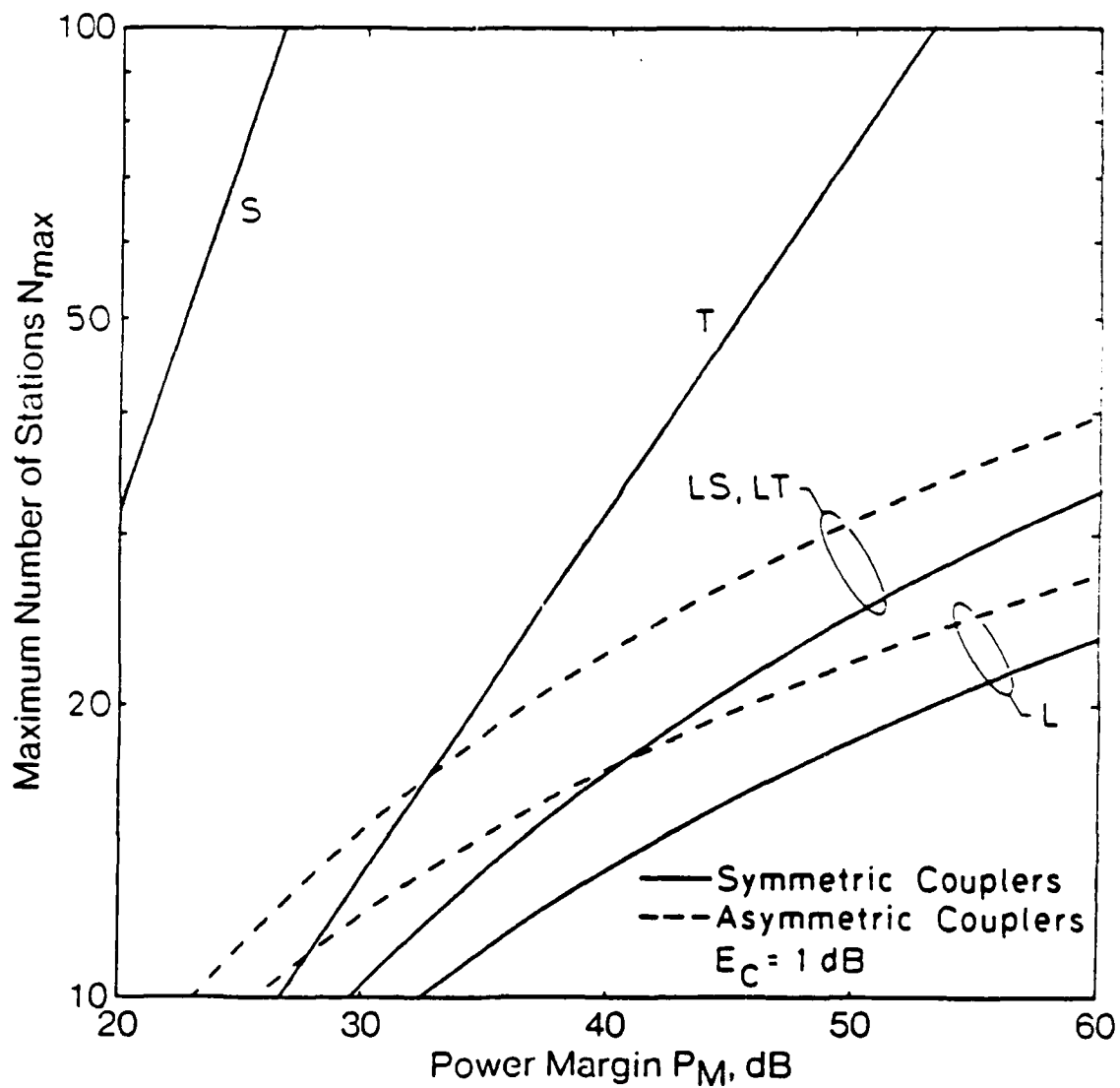


Fig. 4.19 Maximum number of stations versus power margin in the basic and hybrid configurations for a coupler excess loss of 1 dB.

couplers, as compared to the corresponding configurations based on asymmetric couplers. Secondly, we note that a significant portion of the disadvantage of the tree configuration over the star configuration is also due to the reciprocity loss*. The reciprocity loss of a coupler in the tree collection subnetwork is $10 \log_{10} 2 (\approx 3)$ dB. Since there are about $\log_2 N$ couplers from every transmitter to the funneling point, the total loss due to reciprocity is $10 \log_{10} N (\approx 3 \log_2 N)$ dB. (In fact, this loss is incurred in any binary tree collection subnetwork, not necessarily having minimum depth, with individually optimized coupling coefficients. See Appendix A.) In the star configuration, of course, no reciprocity loss is present.

The second factor affecting the performance of a configuration is the excess loss. In the star configuration the excess loss between the transmitters and receivers is equal to that of the star coupler, namely $E_C \log_2 N$. In the tree configurations the number of couplers between every transmitter and receiver is equal to the number of levels in the collection and distribution trees, or $\log_2 N$. The total excess loss is therefore equal to $2E_C \log_2 N$, or twice that of the star configuration. In the linear configuration the worst-case excess loss is between the transmitter of the first station and the receiver of the last station and is equal to $2NE_C$. (Recall that no bypass is used here; for large N , bypass reduces this loss.) The worst-case excess loss in both linear-star and linear-tree hybrid configurations is equal to $E_C(N + \log_2 N)$.

In the compound configurations, as S increases the loss due to reciprocity is reduced; hence more stations can be supported. However, this is achieved at the cost of longer fiber. Therefore it is desirable to know how one can trade fiber length for maximum number of stations. We assume that the stations are equally spaced along a straight line and normalize the total fiber length to twice the distance

*Although not considered here, non-reciprocal couplers could be used at few levels of the tree collection network; however, since for each of these couplers an output with a larger core diameter is used, the use of non-reciprocal ports at many levels of the tree leads to impractically large core diameters near the funneling point.

between the first and the last stations. Based on these assumptions, the normalized fiber length for the star configuration is exactly N (see Fig. 4.1).

The trade-off between the maximum number of stations and the normalized fiber length, for an excess loss of 1 dB per coupler and a power margin of 40 dB, is shown in Fig. 4.20. The curves in this figure are generated by increasing S , starting from 1. (See the Appendix for the formulas used for the fiber length.) As S is increased both the reciprocity and excess losses decrease, so that a larger number of stations may be supported. As it can be observed, the relationship between the number of stations and the fiber length is almost linear.

4.5 Summary

A number of fiber optic control and data configurations for the attempt-and-defer DAMA schemes were proposed. A general and unified approach to the power budget analysis and optimization problem was presented. Based on this approach, it was shown that in the control configurations the use of stretch and bypass techniques may significantly increase the maximum number of stations. A set of data configurations based on linear, star, and tree topologies were proposed and evaluated. The performance improvements achieved by the use of non-reciprocal couplers were also shown.

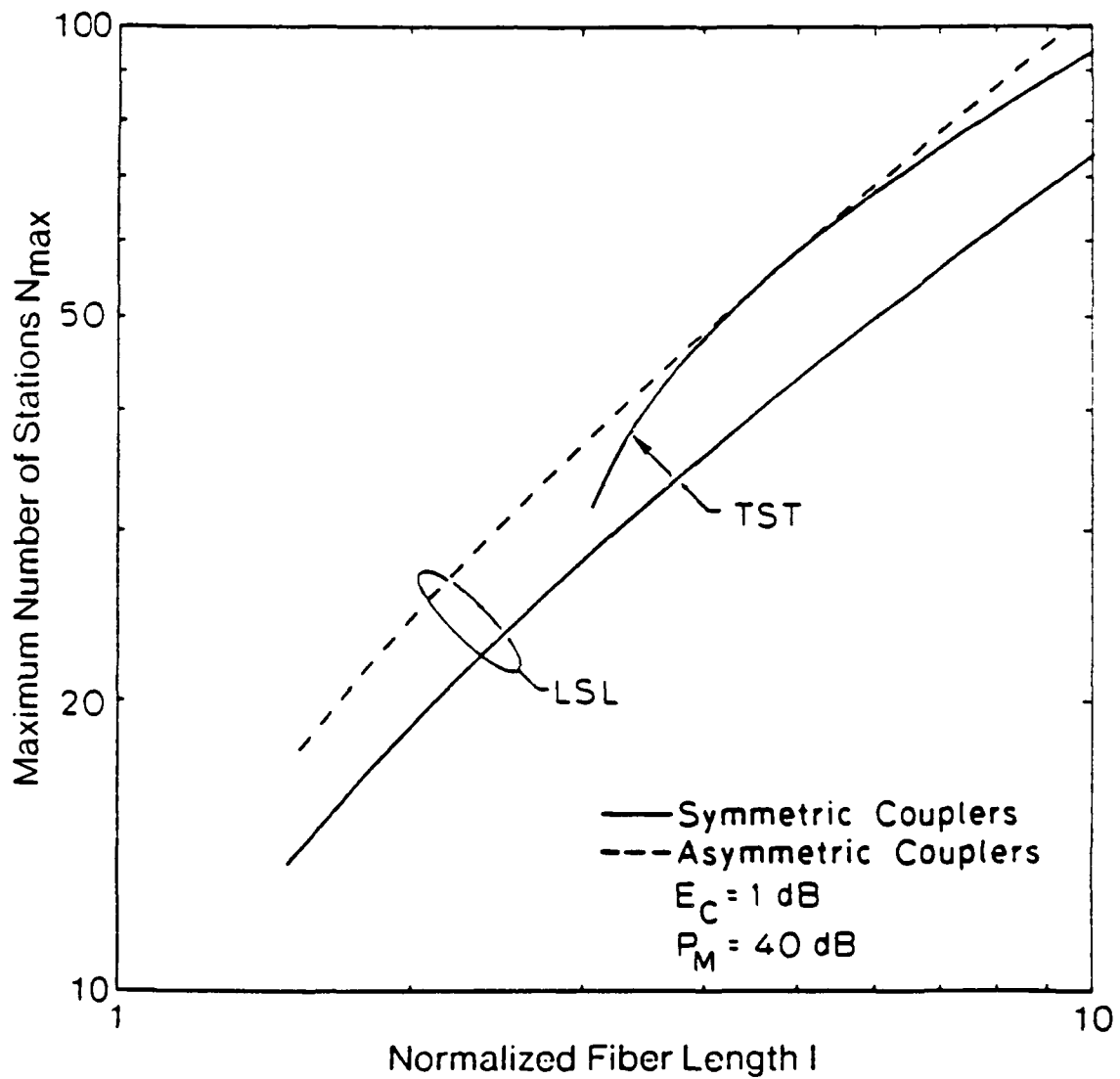


Fig. 4.20 Trade-off between maximum number of stations and fiber length in the compound configurations for a coupler excess loss of 1 dB and a power margin of 40 dB.

Chapter 5

Conclusions and Future Research

5.1 Conclusions

In this work, we addressed the problem of integrated services on local area networks. We indicated that, due to the multitude of applications contemplated in local environments, the volume of traffic expected on local area networks will undoubtedly require high speed operation of such networks; accordingly the solution to integrated services local area networks will more naturally be found in unidirectional networks with attempt-and-defer access mechanism, such as Expressnet and Fasnet, as these are more efficient than the existing standard networks. Furthermore, it was indicated that the implementations of the physical layer of high speed unidirectional LAN's will have to be based on fiber optics technology due to the high bandwidth characteristic of such a medium.

The contributions of this work fall in two areas: (i) the incorporation of transmission scheduling policies into the design of unidirectional local area networks in view of supporting applications with differing requirements, and (ii) The study of a

number of fiber optics topological configurations appropriate for the implementation of unidirectional Local Area Networks.

In chapter 2, we examined a number of transmission scheduling policies under ideal conditions (i.e., zero propagation delay, and complete knowledge of all information regarding messages to be scheduled). Under such conditions, the network could be viewed as a single-server queueing system. A delay-cost function was assigned to each message, and the problem consisted of determining a scheduling policy which minimizes the expected cost per message. As the optimization problem in general is quite complex, we considered six policies and compared their performance. While the first four (first-come-first served (FCFS), round-robin (RR), highest-penalty-first-served (HPFS), and lowest-slack-first-served (LSFS)) had already existed in the literature on scheduling policies, the other two (earliest-slot-first-served (ESFS), and dynamic priority (DP)) constitute new proposals. Simulation was used to put the six policies to the test. Both step and ramp delay-cost functions were considered. An arbitrary and general traffic model was used, in which the delay allowance and cost penalty for each message are randomly selected. All results considered have shown the dynamic priority policy is superior to all others and thus is suboptimal.

In chapter 3, a network structure was proposed which can accommodate the transmission scheduling policies discussed. (In this chapter, the assumptions leading to ideal condition have been relaxed.) We chose to consider a distributed implementation for two reasons: reliability and improved performance. We proposed (i) a broadcast mechanism for dissemination of all relevant information about outstanding messages as required by the the scheduling policies, and (ii) a station architecture which allows the implementations of any of the scheduling policies described in chapter 2.

The broadcast scheme was based on the attempt-and-defer mechanism over a

folded bus structure, such as expressnet. This provided necessary features to guarantee consistency in the resulting decision made by stations regarding the message to be transmitted at a given time.

In chapter 4, we investigated fiber optics topological configurations for unidirectional integrated-services local area networks. We first extracted the special features of doubly folded structure which proved particularly useful in devising fiber optic configurations. We identified two separable functions: (i) the linear ordering among the stations which is essential for an attempt and defer access mechanism, and (ii) data transmission. We have pointed out that these two functions may be provided by separate control and data subnetworks, as this could lead to a higher number of stations that can be accommodated in the network. We have also identified the separability of collection subnetwork and the distribution subnetwork in the data transmission networks, and took advantage of this fact in devising the various configurations.

For the linear control subnetwork, we have proposed two novel topologies: the stretched control configuration and the bypass control configuration. For the overall network implementation, we have grouped the many proposed configuration into three classes: basic, hybrid, and compound.

In chapter 4, we have also undertaken a complete power budget analysis of the fiber optic configurations. In particular, we have provided a general and unified approach to the power budget analysis, and to the coupler optimization problem, and applied the technique to the study of the fiber optic configurations. The results were given in terms of maximum number of stations that each configuration could accommodate for a given power margin. (We have also derived simple (but loose) bounds on the number of stations.) It was shown that the stretched configuration for control subnetwork is superior to the straightforward implementations

using two couplers per stations, and that for power margin above 35dB the bypass configuration is superior to the stretched control configuration. As far as the data transmission configurations are concerned, the star topology is the best; the next is the tree configuration; below these come the hybrid configurations (i.e., linear/star and linear/tree); the worst performance is attained with all linear topology.

Finally, in Appendix A, we provided further results on the performance of fiber optic configurations. In particular, we have taken into account the effect of component parameters explicitly, and studied the sensitivity of performance to these parameters. In the appendix, we have also considered the uniform optimization of coupling fraction; i.e., the optimization under the constraint that all coupling fractions be equal.

The numerical results have shown that, with multimode fibers, high-cost components operating at 100Mps, and uniform optimization of couplers, all fiber-optics configurations implementing folded bus structure can support more than 50 stations, with the exception of the all-linear configuration (i.e., linear collection and linear distribution) which can support 30 stations. For configurations in which the *data (collection) subnetworks* do not have a linear topology (i.e., tree, star, or compound), the number of stations that these data subnetworks can support (discarding the control subnetwork) far exceed 50, reaching several hundreds. The limitation of 50 to 70 stations, depending on the data rate (50 Mbps to 200 Mbps,) is then imposed by the linear control subnetwork

5.2 Suggestions for Additional Research

Several areas of related work have been untouched in this thesis. First of all, the performance evaluation of the scheduling policies has been conducted based on a

general (but arbitrary) model of message traffic and cost functions. It is important to determine means by which to select the appropriate cost functions for messages belonging to different applications in view of attaining the desired effects, and to test the approach proposed in this thesis (i.e., based on transmission scheduling policies,) in realistic scenarios of services.

In the transmission scheduling problem, which was formulated in chapter 2, each cost function was defined on the delay of a single message. However, in some applications, it is important to have some correlation among the message delays. For example, consider voice traffic where message loss of upto 5% is tolerable. If the cost functions are defined on the delay of single messages, the overall probability of loss less than 5% can be achieved. But there is a possibility that lost messages are lumped together, resulting in a significant degradation in voice quality. To avoid this problem, the cost function must be defined on the delays of successive messages. For example, for the voice traffic, one may define the cost functions on the delays of a block of messages. Then for a block size of say 20, the cost of losing more than one message will be made very high. This will increase the separation between the lost messages to acceptable level. Defining the delay-cost function on a block of messages makes the transmission scheduling problem more complex than when delay-cost functions are defined on single messages. Since the original transmission scheduling was analytically intractable, the new and more general problem is also analytically intractable. Therefore, one must seek suboptimal solutions and resort to simulation to test their performance.

There are two other areas which were not covered in this thesis. The first is the VLSI implementation of the station's architecture. Such an implementation must be simple and reliable. It must also have short latencies, such that the overall performance is not degraded significantly. The other area of research is on the fiber

optic configurations. All configurations presented in chapter 4 were based on passive bus topologies. However, to increase the number of stations supported on a bus, it may be necessary to introduce active elements, such as regenerative repeaters. This will increase the number of stations, but at the same time will reduce the reliability of the network, since active elements in general fail much more frequently than the passive elements do. It is important to quantify the tradeoff between between reliability and number of stations.

Appendix A

Additional Results on Fiber Optic Configurations

In a general network, the coupling fractions resulting from maximization of transmittance are not necessarily the same. But, in order to facilitate the maintainability and flexibility of the network, we may need some degree of uniformity among these coupling fractions. When there are no uniformity constraints, we refer to the problem of maximizing the transmittance as *Individual Optimization of Couplers* (IOC), otherwise we call it *Uniform Optimization of Couplers* (UOC). In this appendix we present the performance of the fiber optic configurations under uniform optimization of couplers and compare it with that under individual optimization of couplers. First, we briefly review the characteristics of the fiber-optics components used in the construction of the transmission medium.

A.1 Fiber-Optics Components

The fiber-optics components needed are the fiber, connectors, splices, couplers, transmitters, and receivers. In describing the characteristics of these components,

throughout this appendix and as in chapter 4, we denote the *transmittance* from a point A on the medium to a downstream point B , by $G(A, B)$. $G(A, B)$ represents the fraction of the power at A which reaches B .

Fibers :

Pulse dispersion and optical loss are the most important considerations which affect the choice of fiber and wavelength of light. Pulse dispersion limits the fiber bandwidth [Li78]. However, for the data rates and network dimensions required in local area communications, multimode as well as singlemode fibers with sufficient bandwidth are readily available, and thus pulse dispersion is of no concern in this study. The optical loss of the fiber, however, will be taken into account in our analysis. Let f denote the fiber attenuation in dB/km. Then, the transmittance from a point A to a point B due to l kilometers of fiber is $10^{-lf/10}$.

Connectors and Splices :

A connector is represented graphically by a circle as shown in Fig. A.1(a). It is characterized by its insertion loss in dB, denoted by L_C . The transmittance across a connector is thus given by $10^{-L_C/10}$. As commercially available connectors may introduce higher loss than desired, in order to reduce the total loss through the network, one may connect the fibers by *splicing* them. Letting L_S denote the insertion loss of a splice in dB, (typically $L_S < L_C$), the transmittance through it is $10^{-L_S/10}$. As the connection of two fibers can be implemented either by a connector or by a splice, we refer to such a connection as a *joint*; graphically a joint is represented by a bullet, as shown in Fig. A.1(b). The insertion losses of a joint in dB is denoted by L_J .

Couplers :

In general, an $m \times n$ coupler is a device with m input ports and n output ports which distributes the power from any of the input ports among all output ports.

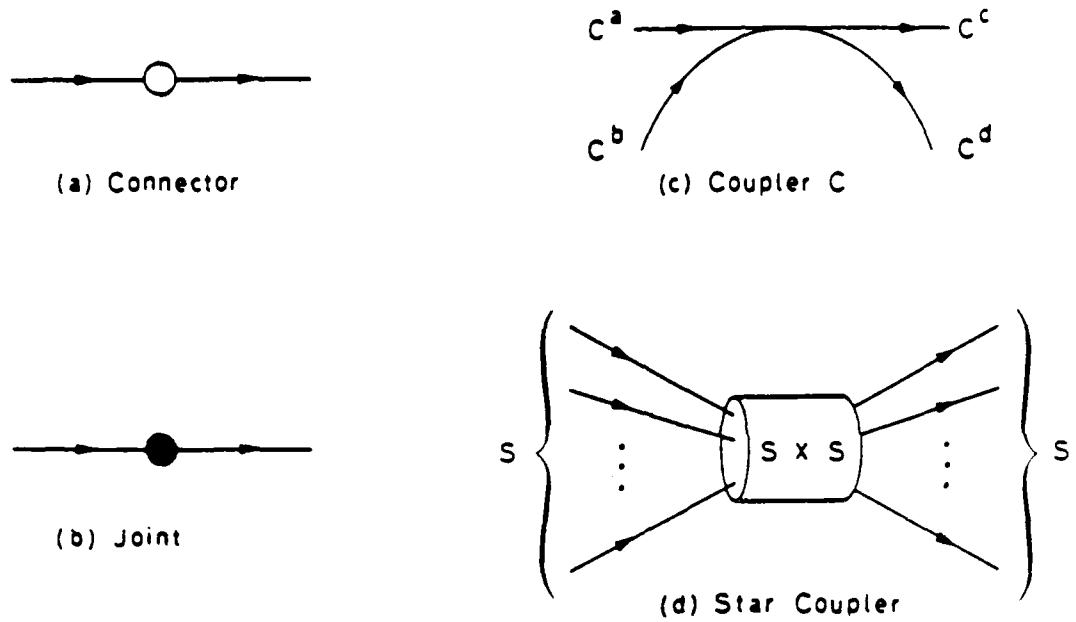


Fig. A.1 Fiber-optics components: (a) a connector, (b) a joint, (c) a coupler, and (d) an $S \times S$ star coupler.

We represent a 2×2 coupler by a directed line segment tangent to a directed arc, as shown in Fig. A.1(c). Since this representation is not symmetric with respect to any pair of the ports, it allows us to identify the ports uniquely. As shown in the figure, we denote the four ports of a coupler C by C^a , C^b , C^c , and C^d . Due to the reciprocity property of optical couplers, $G(C^a, C^d)$ is equal to $G(C^b, C^c)$, and $G(C^a, C^c)$ is equal to $G(C^b, C^d)$. Let E_C denote the excess-loss in dB through a coupler and α denote the corresponding transmittance, i.e. $\alpha \triangleq 10^{-E_C/10}$. Then, for some coupling fraction x , the transmittances through the coupler C are given by,

$$G(C^a, C^d) = G(C^b, C^c) = \alpha x \quad (A.1)$$

$$G(C^a, C^c) = G(C^b, C^d) = \alpha(1 - x)$$

Depending on the way the couplers are implemented, the range of coupling fractions may be more limited than that from 0 to 1. Let $x_{\min}$ and $x_{\max}$ denote the upper

S	$10 \log_{10} S + E(S)$	$E(S)$	$E(S)/\log_2 S$
4	8.0	2.0	1.0
8	11.5	2.5	0.83
16	16.0	4.0	1.0
32	19.0	4.0	0.8
64	23.0	5.0	0.83

Table A.1 The excess-loss for a family of $S \times S$ star couplers [CANSTAR]

and lower limits on the coupling fractions. Normally, we have $0 \leq x_{\min} < 0.5 < x_{\max} \leq 1$.

A *Star* coupler is an $S \times S$ coupler which ideally distributes the power received at any of its input ports equally among its output ports. A star coupler is depicted in Fig. A.1(d). Let $E(S)$ denote the excess-loss in dB through an $S \times S$ star coupler. Then the transmittance through the coupler is $\frac{10^{-E(S)/10}}{S}$.

The excess-loss for a family of star couplers is listed in table 1 [CANSTAR]. Note that, as a good approximation, the excess-loss grows as $\log_2 S$. For a modular star coupler based on 2×2 couplers [MARH84], the excess-loss is approximately $E_C \log_2 S$. Therefore, in our computations, $E(S)$ is assumed to be equal to $E_C \log_2 S$.

Transmitters and Receivers :

An optical transmitter is a system consisting of a light source, a means for efficiently coupling the source's output power into the transmission fiber, and a modulation circuit. In the case of a laser, the system also includes a level control circuit. Transmitters are characterized by the maximum power they can inject into the fiber. In terms of modulation rates, and to a certain degree the output power, a typical LED system is quite limited compared to a laser system. However, for data rates on the order of 50 Mbps, LED systems with reasonable output power (about 1 mW) are commercially available [MIDW83].

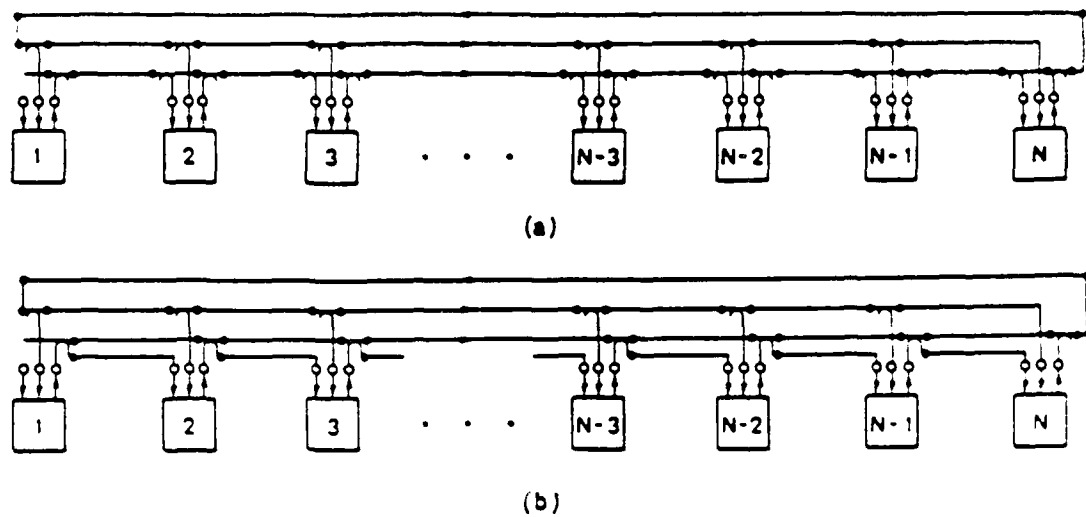


Fig. A.2 Linear configuration (L).

An optical receiver consists of an optical detector and a demodulation circuit. A receiver is characterized by its *sensitivity*; i.e., the minimum received power it requires at given data and error rates. Due to several physical and cost constraints, the sensitivity of optical receivers is higher than that of metallic cable receivers [RHOD83]. Furthermore, for a given bit error rate, the higher the data rate provided by a receiver is, the lower is its sensitivity [OGAW83]. In an actual implementation, since the sensors are not required to detect individual bits, their sensitivity may be somewhat better than that of data receivers. In section 4.4, we discuss the implications of the higher sensitivity of sensors.

A.2 Description of Configurations

As we have indicated in chapter 4, there are many fiber-optics configurations to implement the folded-bus structure. These have been grouped into three classes: *basic*, *hybrid*, and *compound*. In figures A.2-A.8, we give an accurate representation

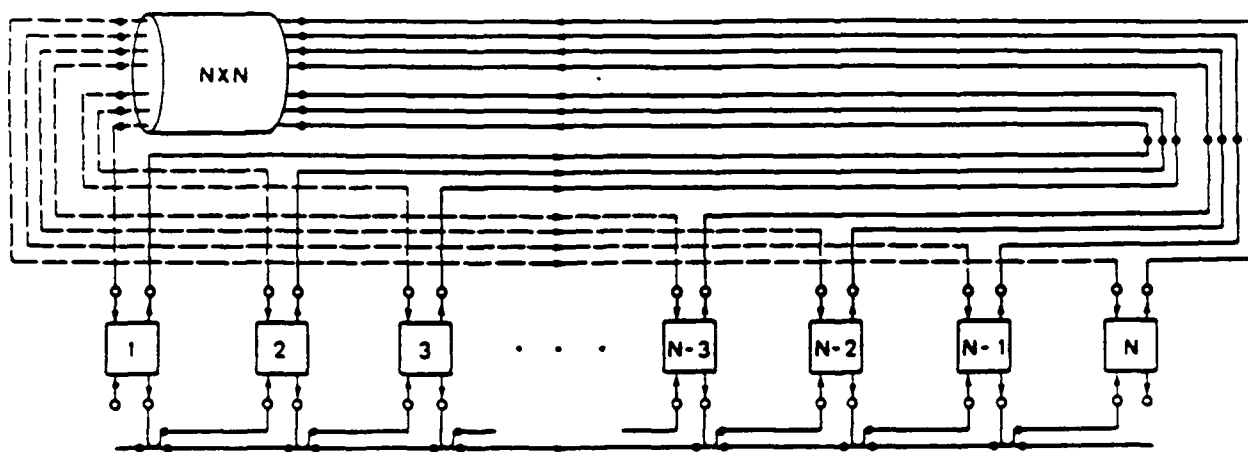


Fig. A.3 Star configuration (S/C).

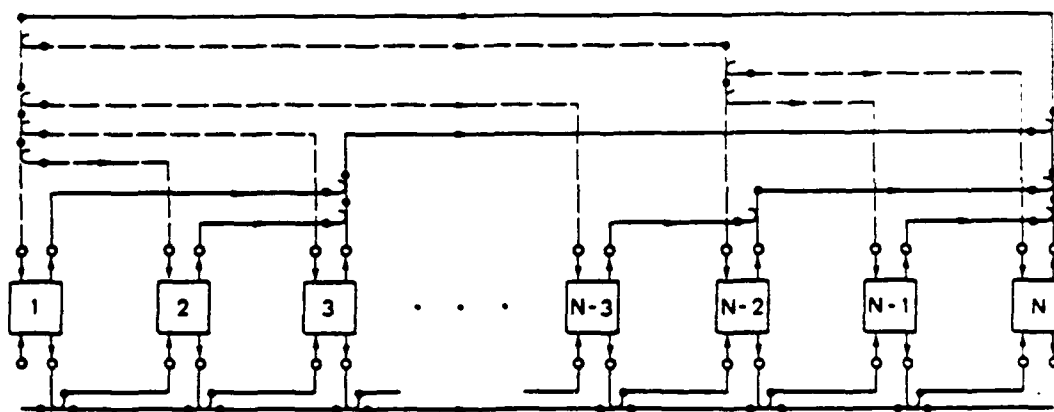


Fig. A.4 Tree configuration (T/C).

of the configurations satisfying the delay constraint and identifying all fiber optic components involved.

A.3 Uniform Optimization of Couplers

In the uniform optimization of couplers, $\underline{G}(S; x)$ is maximized under a set of

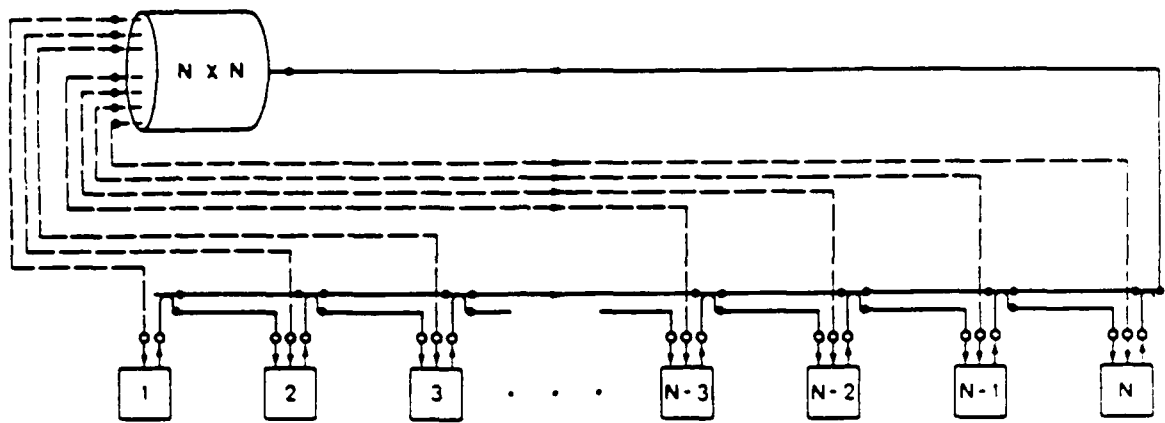


Fig. A.5 Linear/star configuration (LS).

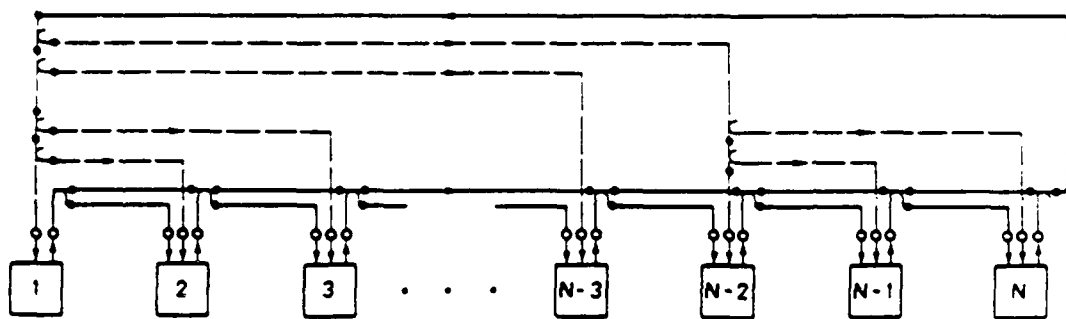


Fig. A.6 Linear/tree configuration (LT).

uniformity conditions on the coupling fractions. The most obvious condition of uniformity is to require that all the coupling fractions in the network be the same. However, since the collection subnetwork, the distribution subnetwork, and the control subnetwork may have entirely different topologies, such a condition may in some instances lead to excessively small $\underline{G}(S; \mathbf{x})$. (For example, this is the case for the LT configuration as will be seen in section A.3.3 below.) Moreover, maintainability and flexibility are not jeopardized if different couplers are used in these

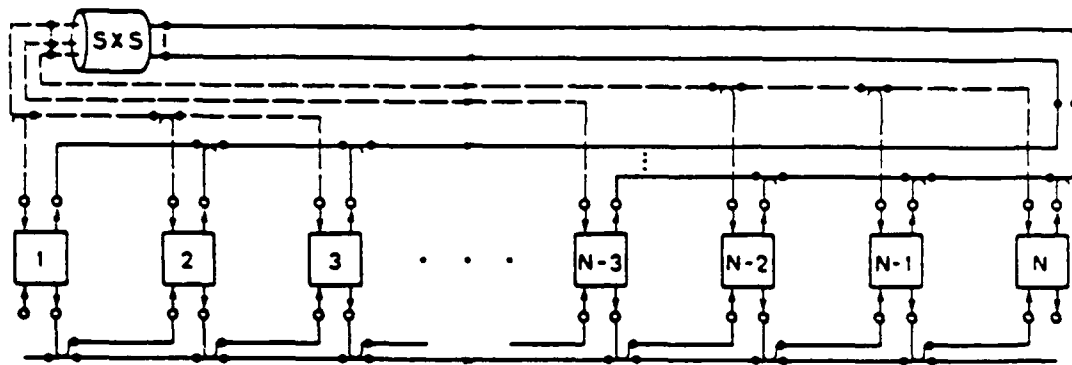


Fig. A.7 Linear/star/linear configuration (LSL/C).

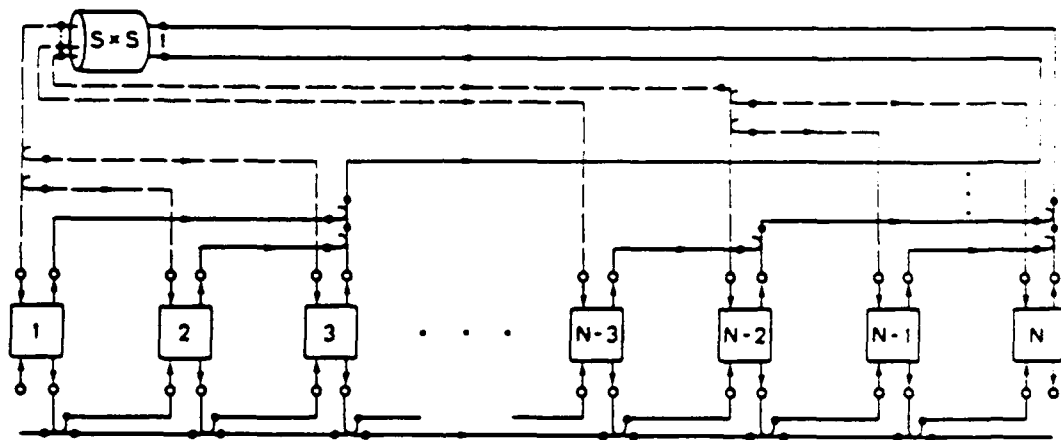


Fig. A.8 Tree/star/tree configuration (TST/C).

subnetworks. Hence we perform the optimization on these subnetworks separately. The subnetworks considered here are: the linear control, linear collection(/control), linear distribution, tree collection, and tree distribution subnetworks.

Let x denote the common coupling fraction in the subnetwork under consideration. If the transmittance from a point A to a downstream point B is a not a function of x , it is denoted by $G(A, B)$; otherwise, it is denoted by $G(A, B; x)$. Note

that $G(A, B; x)$ can be separated into two factors. The first factor is constant with respect to x and is denoted by $G_0(A, B)$. The second factor is $x^m(1-x)^n$, for some nonnegative intergers m and n such $m+n \geq 1$. For station $i = 1, 2, \dots, N$, let T_i , R_i , and R'_i denote its transmitter, receiver, and sensor, respectively. Here, for the sake of simplicity we assume that all transmitters have the same output power and all receivers have the same sensitivity.

A. Linear Control Subnetwork:

Consider the control bus which is depicted in Fig. A.9(a). Our objective is to maximize

$$\underline{G}(x) \triangleq \min_{1 \leq i < j \leq N} G(T_i, R'_j; x) \quad x_{\min} \leq x \leq x_{\max} \quad (A.2)$$

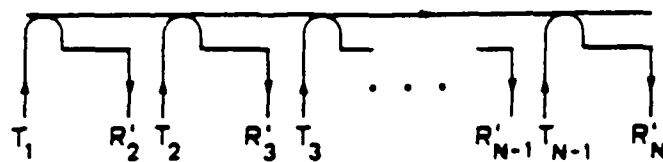
From Fig. A.9(a), we have

$$G(T_i, R'_j; x) = G_0(T_i, R'_j) \times \begin{cases} 1-x & 2 \leq i+1=j \leq N \\ x^2(1-x)^{j-i-2} & 2 \leq i+1 < j \leq N \end{cases} \quad (A.3)$$

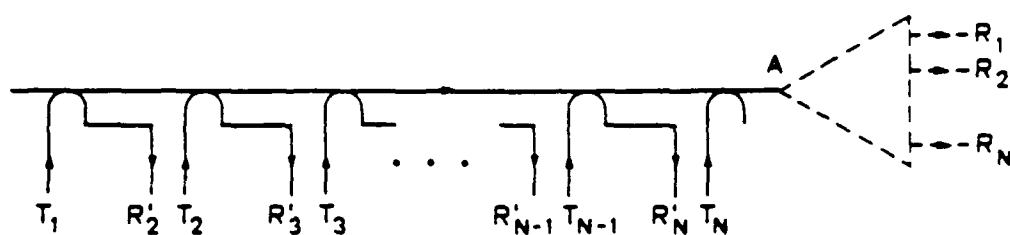
$$G_0(T_1, R'_N) \leq G_0(T_i, R'_j) \quad 1 \leq i < j \leq N$$

For $N = 2$, the optimum x is clearly $x_{\min}$. For $N = 3$, there are three transmitances under consideration, the minimum of which is $\underline{G}(x)$. These are plotted in Fig. A.10. From this figure, we observe that if the optimum x is strictly inside the interval $(x_{\min}, x_{\max})$ then it must be equal to the positive solution of the following quadratic equation.

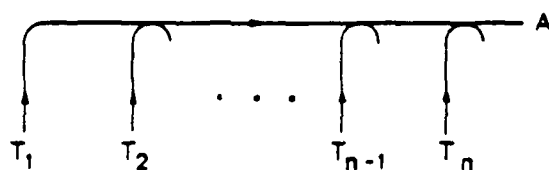
$$G_0(T_1, R'_3)x^2 = \min\{G_0(T_1, R'_2), G_0(T_2, R'_3)\}(1-x) \quad (A.4)$$



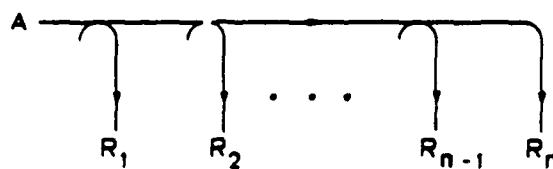
(a) Linear Control Subnetwork



(b) Linear Collection/Control Subnetwork

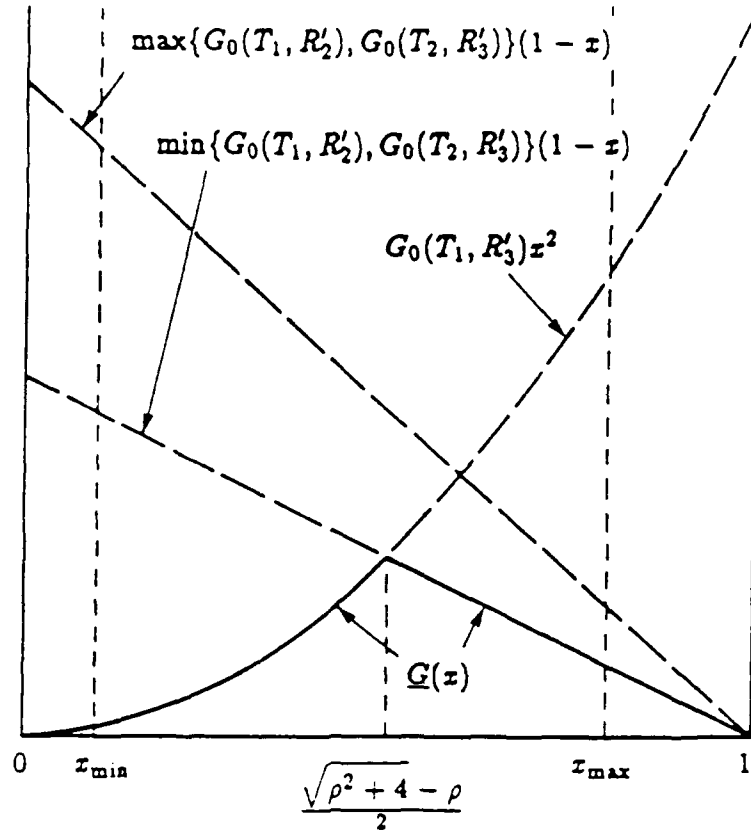


(c) Linear Collection Subnetwork



(d) Linear Distribution Subnetwork

Fig. A.9 The subnetworks on which the uniform optimization is performed.



Uniform Coupling Fraction x

Fig. A.10 $\underline{G}(S; x)$ as a function of x for linear control subnetwork and $N = 3$.

Let

$$\rho \triangleq \frac{\min\{G_0(T_1, R'_2), G_0(T_2, R'_3)\}}{G_0(T_1, R'_3)} \quad (A.5)$$

Then the optimum x is equal to $\left[\frac{\sqrt{\rho^2 + 4\rho} - \rho}{2} \right]_{x_{\min}}^{x_{\max}}$. For $N \geq 4$, a more interesting case, we have a simpler result due to the fact that

$$\underline{G}(x) = G(T_1, R'_N; x) = G_0(T_1, R'_N)x^2(1-x)^{N-3} \quad (A.6)$$

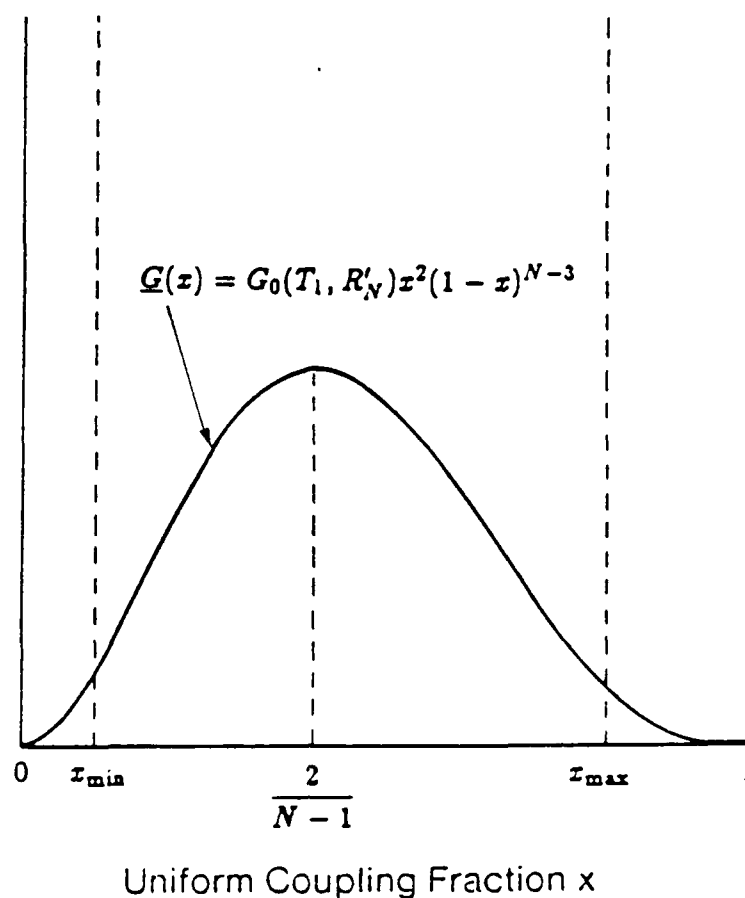


Fig. A.11 $\underline{G}(\mathcal{S}; x)$ as a function of x for linear control subnetwork and $N \geq 4$.

The sketch of function $\underline{G}(\mathcal{S}; x)$ is shown in Fig. A.11. By differentiating the right hand side of (A.6), we find that the optimum x is given by $\left[\frac{2}{N-1}\right]_{x_{\min}}^{x_{\max}}$.

B. Linear Collection/Control Subnetwork:

Consider the linear collection/control subnetwork which is shown in Fig. A.9(b).

Our objective is to maximize

$$\underline{G}(x) \triangleq \min \left\{ \min_{1 \leq i < j \leq N} G(T_i, R'_j; x), \min_{1 \leq i \leq N, 1 \leq k \leq N} G(T_i, R_k; x) \right\} \quad x_{\min} \leq x \leq x_{\max} \quad (A.7)$$

Using Fig. A.9(b), we have in addition to (A.3),

$$G(T_i, R_j; x) = G_0(T_i, A)x(1-x)^{N-i}G(A, R_j) \quad i, j = 1, 2, \dots, N \quad (A.8)$$

$$G_0(T_1, A) \leq G_0(T_1, R'_N)$$

Clearly

$$\min_{1 \leq i \leq N, 1 \leq k \leq N} G(T_i, R_k; x) = G_0(T_1, A)x(1-x)^{N-1} \min_{1 \leq k \leq N} G(A, R_k) \quad (A.9)$$

By differentiating, we observe that the right hand side of (A.9) is maximized at $\left[\frac{1}{N}\right]_{x_{\min}}^{x_{\max}}$. Now, we show that $\left[\frac{1}{N}\right]_{x_{\min}}^{x_{\max}}$ maximizes $\underline{G}(x)$, by proving

$$\min_{1 \leq i < j \leq N} G(T_i, R'_j; \left[\frac{1}{N}\right]_{x_{\min}}^{x_{\max}}) \geq \min_{1 \leq i \leq N, 1 \leq k \leq N} G(T_i, R_k; \left[\frac{1}{N}\right]_{x_{\min}}^{x_{\max}}) \quad (A.10)$$

Recall that $G(A, R_k)$ is the fraction of power at A which reaches R_k . Therefore, we have $\sum_{1 \leq k \leq N} G(A, R_k) \leq 1$ and thus $\min_{1 \leq k \leq N} G(A, R_k) \leq \frac{1}{N}$. From this inequality, together with (A.3) and (A.9), (A.10) follows.

C. Linear Data Collection and Distribution Subnetworks:

Consider the linear data collection subnetwork shown in Fig. A.9(c) with transmitters $T_1, T_2, \dots, T_n$. This in fact may represent one of the collection subnetworks

in a compound configuration. (A distribution subnetwork as shown in Fig. A.9(d) can be treated in a similar way.) Our objective is to maximize

$$\underline{G}(x) \triangleq \min_{1 \leq i \leq n} G(T_i, A; x) \quad x_{\min} \leq x \leq x_{\max} \quad (\text{A.11})$$

From the figure, we have

$$G(T_i, A; x) = G_0(T_i, A) \times \begin{cases} (1-x)^{n-1} & i = 1 \\ x(1-x)^{n-i} & 2 \leq i \leq n \end{cases} \quad (\text{A.12})$$

$$G_0(T_i, A) \leq G_0(T_{i+1}, A) \quad 1 \leq i \leq n-1$$

From (A.12) we note that

$$\underline{G}(x) = \min \{ G_0(T_1, A)(1-x)^{n-1}; \quad G_0(T_2, A) \cdot x(1-x)^{n-1} \} \quad (\text{A.13})$$

From Fig. A.12, we observe that $\underline{G}(x)$ is maximized at

$$x = \left[\min \left(\frac{1}{n-1}, \frac{G_0(T_1, A)}{G_0(T_1, A) + G_0(T_2, A)} \right) \right]_{x_{\min}}^{x_{\max}} \quad (\text{A.14})$$

D. Tree Collection and Distribution Subnetworks:

In the case of the tree subnetworks, an analysis similar to the above becomes intractable; a numerical search method may then have to be used. However, based on the construction of the minimum-depth tree, it should be intuitively clear that the optimum coupling fraction must be very close to $\frac{1}{2}$. (This indeed is the case for all the numerical examples considered in section A.3.3 below.)

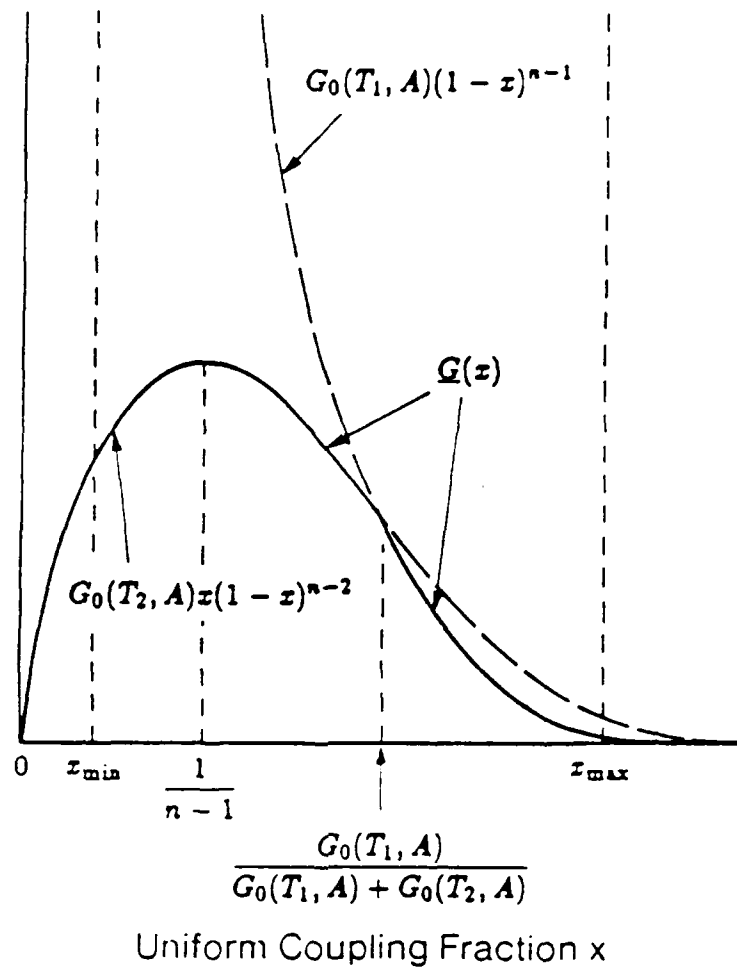


Fig. A.12 $\underline{G}(S; x)$ as a function of x for linear data collection subnetwork.

A.4 Numerical Results

Based on the parameters of commercially available components, we evaluate the maximum number of stations $N_{\max}$ that each of the configurations which was presented in section A.1 can support. The effect of coupling fractions on $N_{\max}$ and the sensitivity of $N_{\max}$ to various parameters are also presented.

A.4.1 Commercially Available Components

We assume that all transmitters have the same output power $\mathcal{P}_T$, and all receivers have the same sensitivity $\mathcal{P}_R$, in Watts. The power margin of the network

Fiber	Parameter	Unit	LC	MR	HC	HC Ref
Multimode or Singlemode	f	dB/km	6	0.98	0.06	[SPECTRUM]
	L_C	dB	1	0.5	0.25	[LASERS]
	L_S	dB	0.3		0.03	[FURUKAWER]
	L_J	dB	1	0.17	0.03	
	P_R	dBm	-20		-45	[PLESSEY]
Multimode	E_C	dB	1	0.5	0.25	[ADC]
	M_C	%	5	2.2	1	[NEC]
	P_T	dBm	-3		17	[SPECTRA]
	P_M	dB	17	32	62	
Singlemode	E_C	dB	1	0.22	0.05	[KOPERA]
	M_C	%	1	0	0	[FIBERNETICS]
	P_T	dBm	2		2	[LASER]
	P_M	dB	22	32	47	

Table A.2 Values used for fiber-optics component parameters. (Receiver sensitivity, P_R , listed is for a data rate of 50 Mbps and a bit-error rate of 10^{-9} .)

in dB is $P_M \triangleq P_T - P_R$, where P_T and P_R are transmitter power and receiver sensitivity in dBm, respectively. (Note that $P_M = 10 \log_{10} \mathcal{M}$.) The range of coupling fractions is assumed to be symmetric about 0.5; i.e., $x_{\max} = 1 - x_{\min}$; denote the minimum coupling fraction in percent by M_C , i.e. $M_C \triangleq x_{\min} \times 100\%$.

The parameter values for a number of commercially available components are listed in table A.2. For every parameter, the value in the High Cost (HC) column corresponds to the best commercially available component, and the value in the Low Cost (LC) column corresponds to a component with a low cost but still reasonably good characteristics. (References for values in the HC column are given in the table.) For every parameter which directly appears in the power budget analysis, an intermediate value equal to the geometric mean of the HC and LC values is given in the Mid-Range (MR) column. Note that for some of the parameters, the values for multimode and singlemode fibers are different and are thus listed separately.

In the case of multimode fiber, the HC and LC values given for the transmitter

power P_T correspond to laser and LED sources, respectively. (A core diameter of $50\text{ }\mu\text{m}$ is assumed.) In the case of singlemode fibers, since an LED cannot inject sufficiently large power, both HC and LC values correspond to a laser source. The values given for receiver sensitivity P_R in table A.2, are for a data rate of 50 Mbps and an bit error rate of 10^{-9} . (Every increase by a factor of two in the data rate corresponds to about 3 dB degradation in the receiver sensitivity. Therefore, at the data rates of 100 and 200 Mbps, the receiver sensitivity is about -42 and -39 dBm, respectively.) Unless otherwise specified, we shall assume that the data rate to be 50 Mbps throughout the remainder of the paper. Also, we assume that the stations are uniformly spaced and the distance between the first and the last stations is one kilometer.

A.4.2 Maximum Number of Stations

The maximum number of stations, N_{max} , for different subnetworks is computed as follows. The power analysis of the previous section is used to evaluate the minimum power margin required to support a given number of stations, N . Using a numerical search method, the maximum N for which the required power margin does not exceed the available power margin is then found. For the values given in table A.2 (in which the data rate considered to be 50 Mbps) N_{max} for all types of subnetworks with the exception of the LSL and TST data subnetworks is displayed in table A.3. (Tables A.4 and A.5 give N_{max} for 100 and 200 Mbps, respectively.) For the the LSL and TST data subnetworks, N_{max} is a function of the funneling width S and is shown in Fig. A.13. (The jaggedness in the curves is due to the fact that as S is incremented by one, the number of stations in each collection or distribution subnetwork can only decrease by an integer.) For these networks, the trade-off between the maximum number of stations and fiber length is shown in Fig. A.14. Note that, for a configuration which has separate data and control

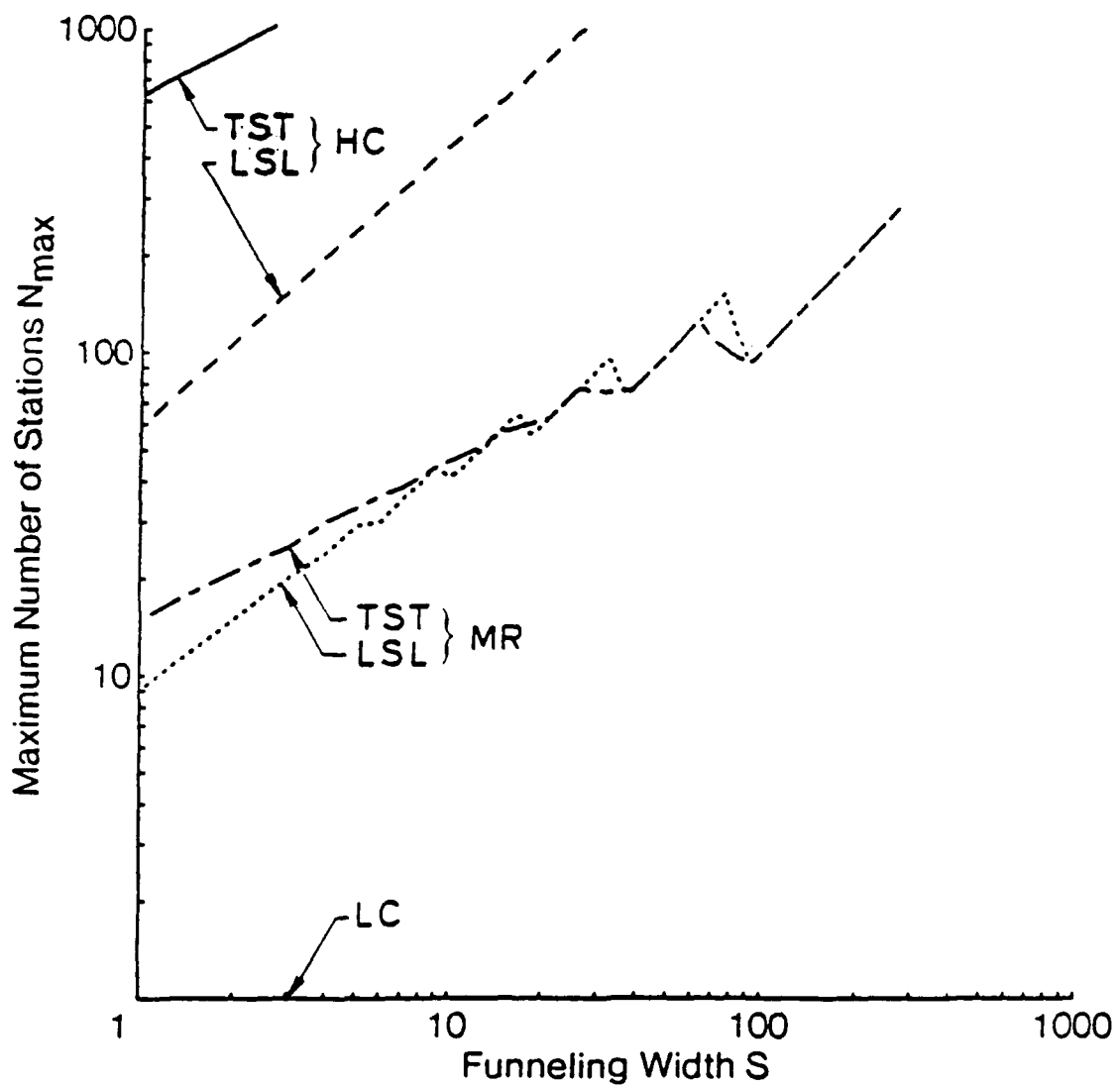


Fig. A.13 The maximum number of stations as a function of the funneling width for multimode fiber and individually optimized couplers.

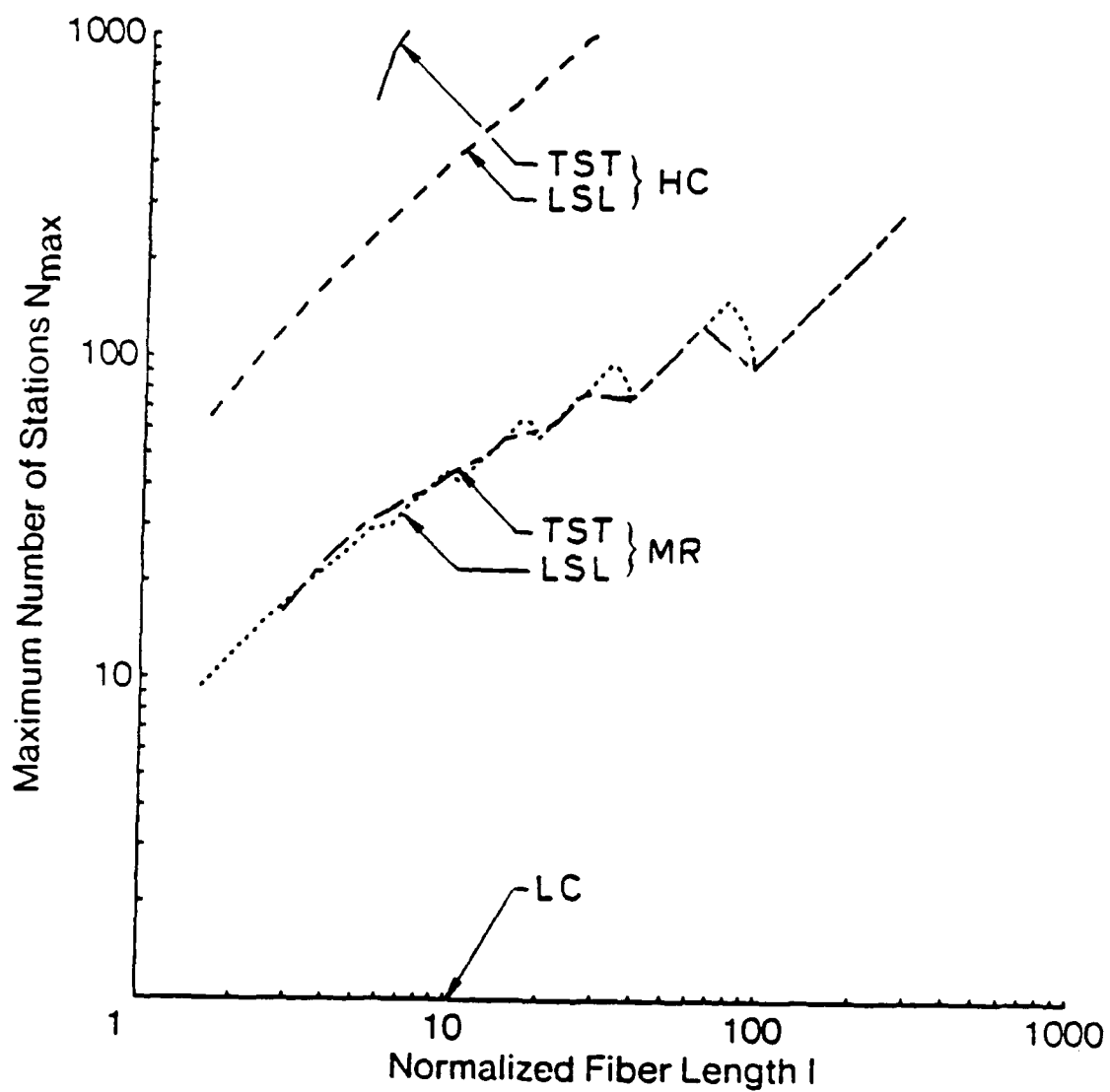


Fig. A.14 The trade-off between the maximum number of stations and fiber length for multimode fiber and individually optimized couplers.

Optimization	Configuration	Multimode Fiber			Singlemode Fiber		
		LC	MR	HC	LC	MR	HC
UOC	C	3	11	72	3	13	69
	L	1	5	35	1	6	34
	T	1	9	512	1	15	128
	S	1	277	461417	1	451	36204
	LT	1	6	62	1	8	54
	LS	1	7	63	1	8	58
IOC	C	4	21	120	5	27	148
	L	1	9	59	1	11	72
	T	1	15	629	1	18	179
	S	1	277	461417	1	451	36204
	LT	1	11	86	1	13	95
	LS	1	11	86	1	14	96

Table A.3 Maximum number of stations for a data rate of 50 Mbps.

subnetworks, $N_{\max}$ is determined by the linear control subnetwork. (assuming, of course, the latter is implemented using fiber-optics technology.)

In generating $N_{\max}$ for subnetworks containing star couplers, no limit has been imposed on the size of the latter. It is to be noted, however, that the largest commercially available star is 100×100 for multimode fibers and 3×3 for single mode fibers. But using a modular approach, larger couplers can be obtained[MARH84].

pp

From the results displayed in table A.3 and Fig. A.13, we note that there are three major factors affecting $N_{\max}$. The first factor is the funneling width: as was discussed in chapter 4, section 4.3.3, the S subnetwork can support significantly more stations than the other data subnetworks. The second factor is the average number of couplers and joints in the paths from transmitters to receivers. The faster this number grows with the number of stations, the smaller is $N_{\max}$. For example, consider the L and T subnetworks. The T subnetwork can support a larger $N_{\max}$

Optimization	Configuration	Multimode Fiber			Singlemode Fiber		
		LC	MR	HC	LC	MR	HC
UOC	C	2	9	65	3	11	57
	L	1	4	32	1	5	28
	T	1	8	256	1	8	67
	S	1	153	243836	1	237	18351
	LT	1	5	54	1	6	40
	LS	1	5	56	1	7	47
IOC	C	3	18	111	4	23	128
	L	1	8	55	1	9	61
	T	1	11	460	1	13	128
	S	1	153	243836	1	237	18351
	LT	1	9	78	1	11	78
	LS	1	9	79	1	11	79

Table A.4 Maximum number of stations for a data rate of 100 Mbps.

Optimization	Configuration	Multimode Fiber			Singlemode Fiber		
		LC	MR	HC	LC	MR	HC
UOC	C	2	7	58	3	9	46
	L	1	3	29	1	4	22
	T	1	5	256	1	8	64
	S	1	69	128855	1	124	9301
	LT	1	4	46	1	4	32
	LS	1	4	50	1	5	38
IOC	C	2	14	103	3	19	105
	L	1	6	50	1	8	51
	T	1	7	335	1	10	91
	S	1	69	128855	1	124	9301
	LT	1	6	71	1	8	62
	LS	1	7	72	1	9	63

Table A.5 Maximum number of stations for a data rate of 200 Mbps.

than the L subnetwork. The average number of couplers on the optical paths in the L subnetwork grows as $\sqrt{N}$ while in the T subnetwork, it grows as $\log_2 N$. The third

factor is the type of optimization of coupling fractions. IOC leads to substantially higher $N_{\max}$ than UOC.

A.4.3 Comments on the Optimum Coupling Fractions

In this section, we examine the numerical values of optimum coupling fractions for various cases, and justify assertions made in section A.2 on uniform optimization.

Fig. A.15 shows the individually optimized coupling fractions in the linear control, linear collection and linear collection/control subnetworks for $N = 20$, multi-mode fiber, and mid-range component parameters. The behavior of the coupling fractions as a function of the station position is very much similar to that for the ideal components, as was discussed in chapter 4, section 4.3.2: in the case of the control subnetwork, the coupling fractions are about 1 for the end stations and decrease as we move toward the center of bus; in the case of the linear collection and collection/control subnetworks, as we move from station 1 to station N the optimum coupling fractions decrease continuously. (Of course, the coupling fraction cannot be less than $x_{\min}$.) Also, by comparing the optimum coupling fractions for the collection subnetwork to those of the collection/control subnetwork obtained by individual optimization, we observe that, although the constraints imposed by sensors are active, their effect on the coupling fractions is not significant. (This is similar to the case of uniform optimization, where we have shown in section A.2 that the sensors' constraints are inactive.)

In section A.2, we stated that $N_{\max}$ may be severely limited if, in uniform optimization, we require the coupling fractions in separate subnetworks to be all equal. An example illustrating this point is the LT data subnetwork for which we have listed, in table A.6, $N_{\max}$ in two cases: case A where we have required the

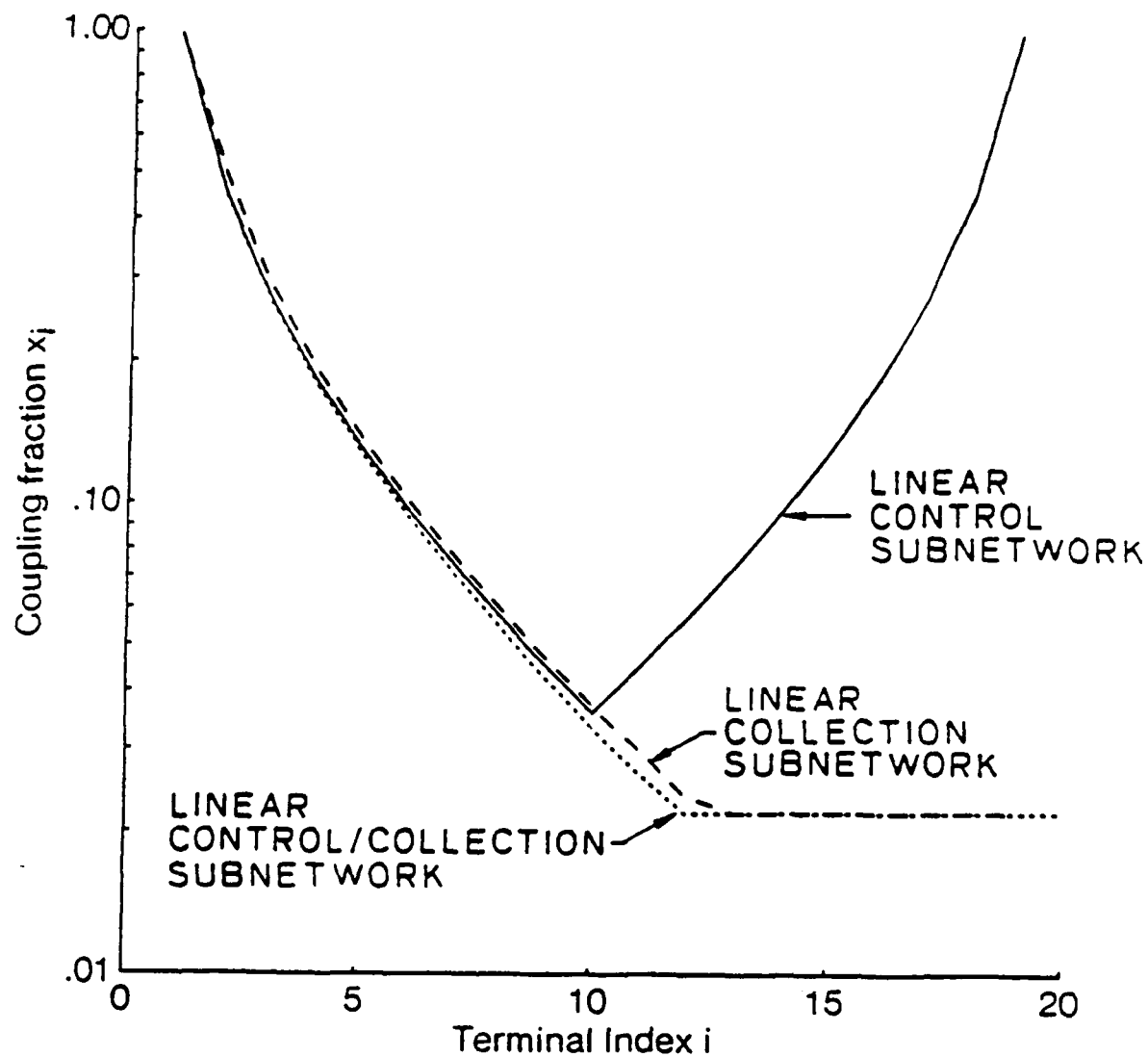


Fig. A.15 Individually optimized coupling fractions in the linear subnetworks for $N = 20$, multimode fiber, and MR components.

Fiber	Components	A	B
Multimode	LC	1	1
	MR	5	6
	HC	23	62
Singlemode	LC	1	1
	MR	5	8
	HC	15	54

Table A.6 $N_{\max}$ in the LT configuration for two cases: case A where the coupling fractions in the collection and distribution subnetworks are required to be the same, and case B where they are not.

coupling fractions in the collection and distribution subnetworks to be the same, and case B where we have not.

In part D of section A.2, we also stated that the optimum value of the coupling fractions in uniform optimization of minimum-depth binary-tree subnetworks is close to $\frac{1}{2}$. We illustrate this fact by displaying in Fig. A.16 $N_{\max}$ as a function of the coupling fraction x , for the T data subnetwork, which peaks at $x = \frac{1}{2}$.

Finally, consider the uniformly optimized coupling fractions in a linear subnetwork, for example C. Fig. A.17 shows the number of stations as a function of the (uniform) coupling fraction x . For the various cases, it is easy to confirm that the optimum x is given by $\frac{2}{N_{\max}+1}$ ($\approx \frac{2}{N_{\max}}$) analytically derived in part A of section A.3. Moreover, Fig. A.17 shows the insensitivity of $N_{\max}$ to small (uniform) variations around the optimum.

A.4.4 Sensitivity of $N_{\max}$ to Component Parameters

The sensitivity of $N_{\max}$ to each component parameter was obtained by varying that parameter over some range while fixing all the other parameters. Based on a large set of numerical results we have made the following observations:

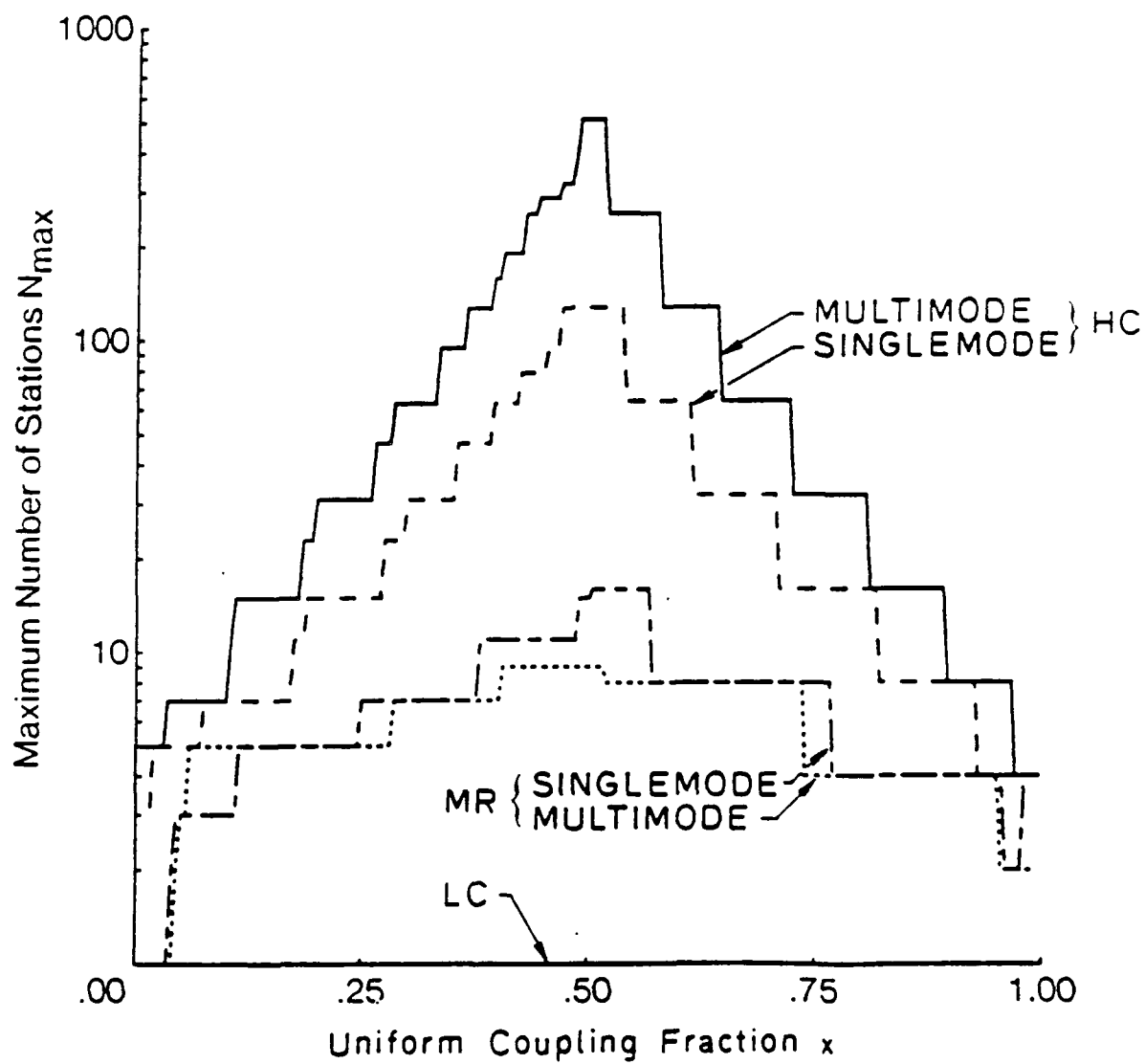


Fig. A.16 N_{max} in the T subnetwork as function of the uniform coupling fraction.

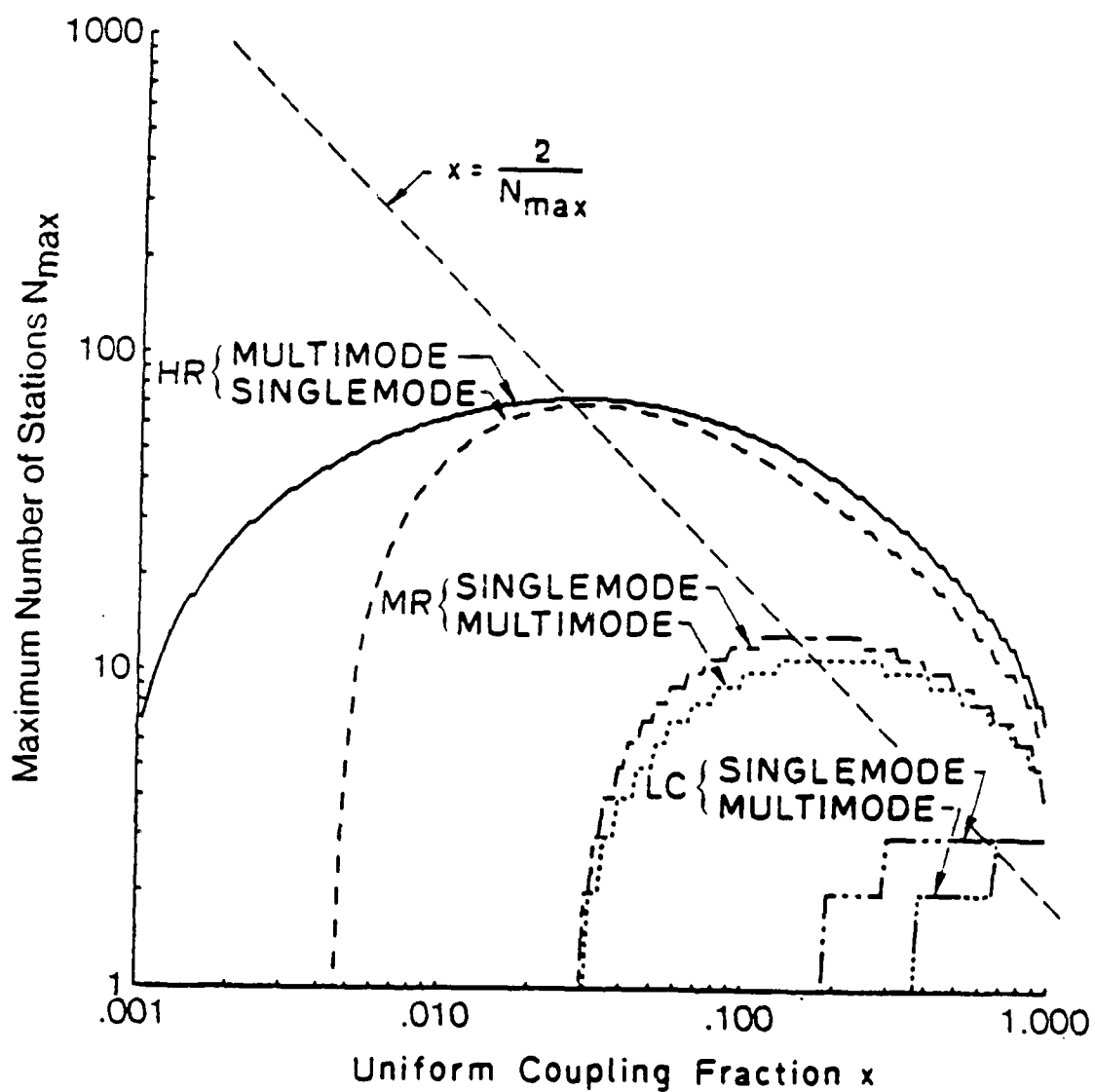


Fig. A.17 $N_{\max}$ in the C subnetwork as a function of the uniform coupling fraction.

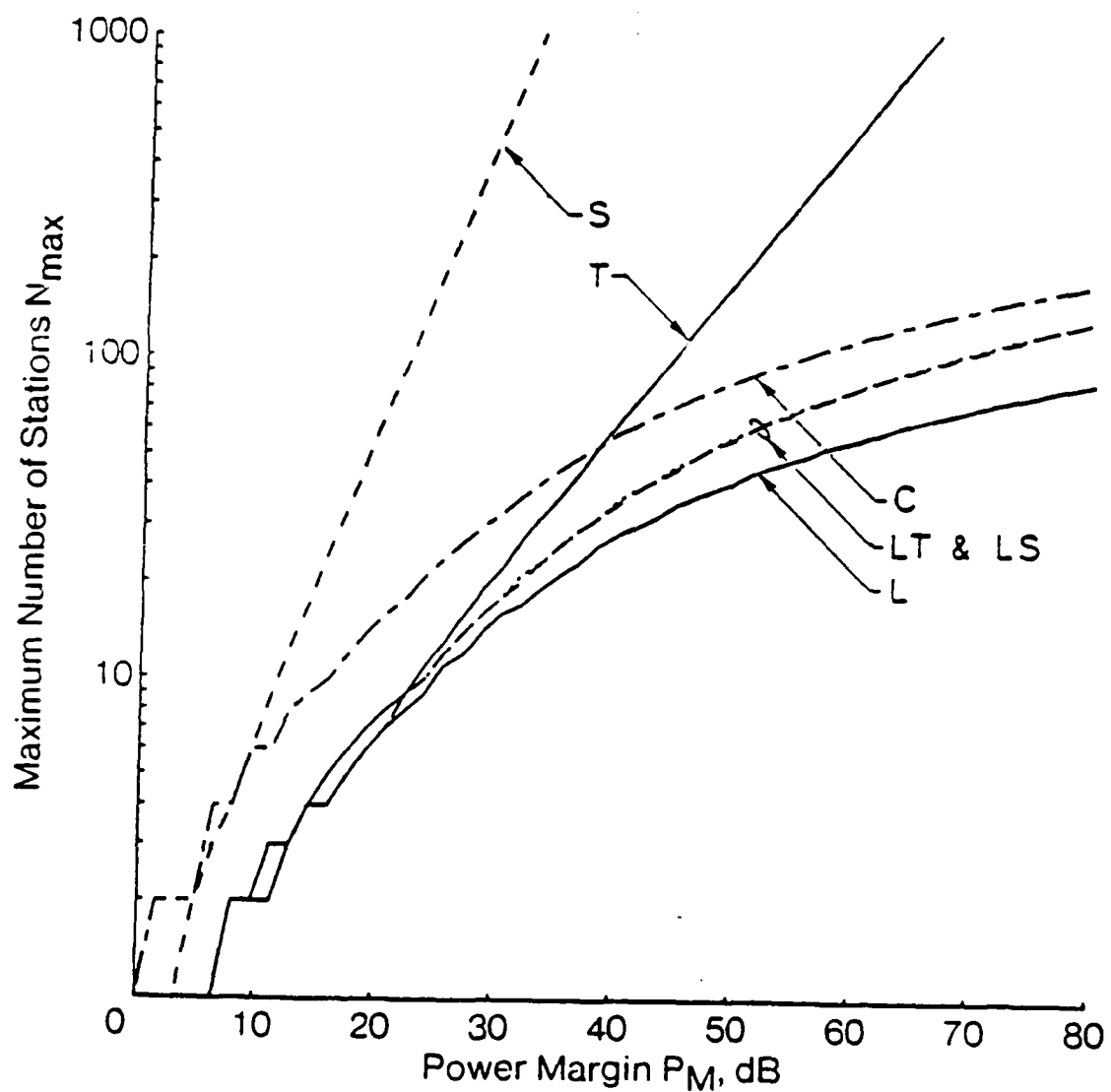


Fig. A.18 The maximum number of stations as a function of the power margin for HC components, multimode fiber, and individually optimized couplers.

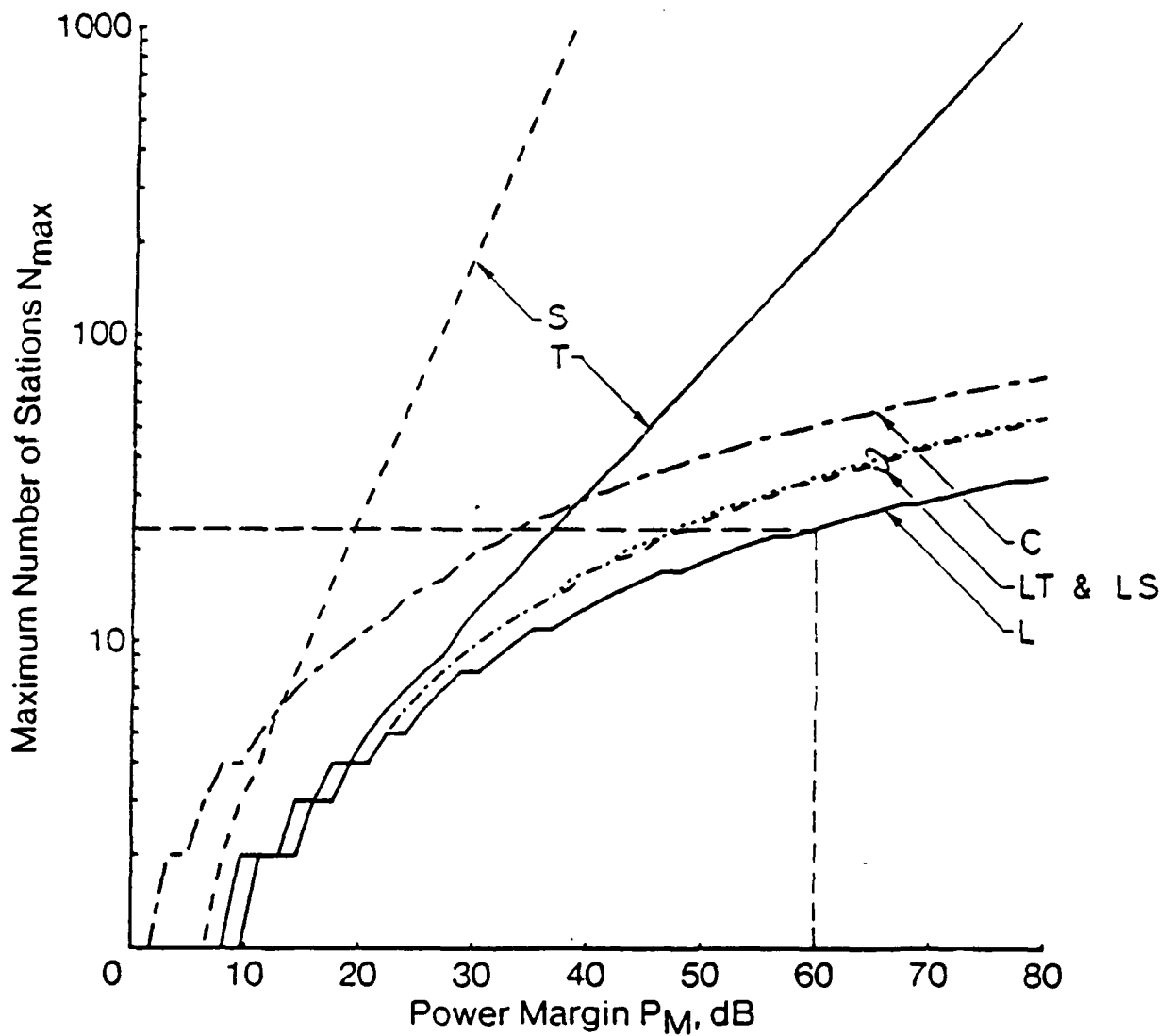


Fig. A.19 The maximum number of stations as a function of the power margin for MR components, multimode fiber, and individually optimized couplers.

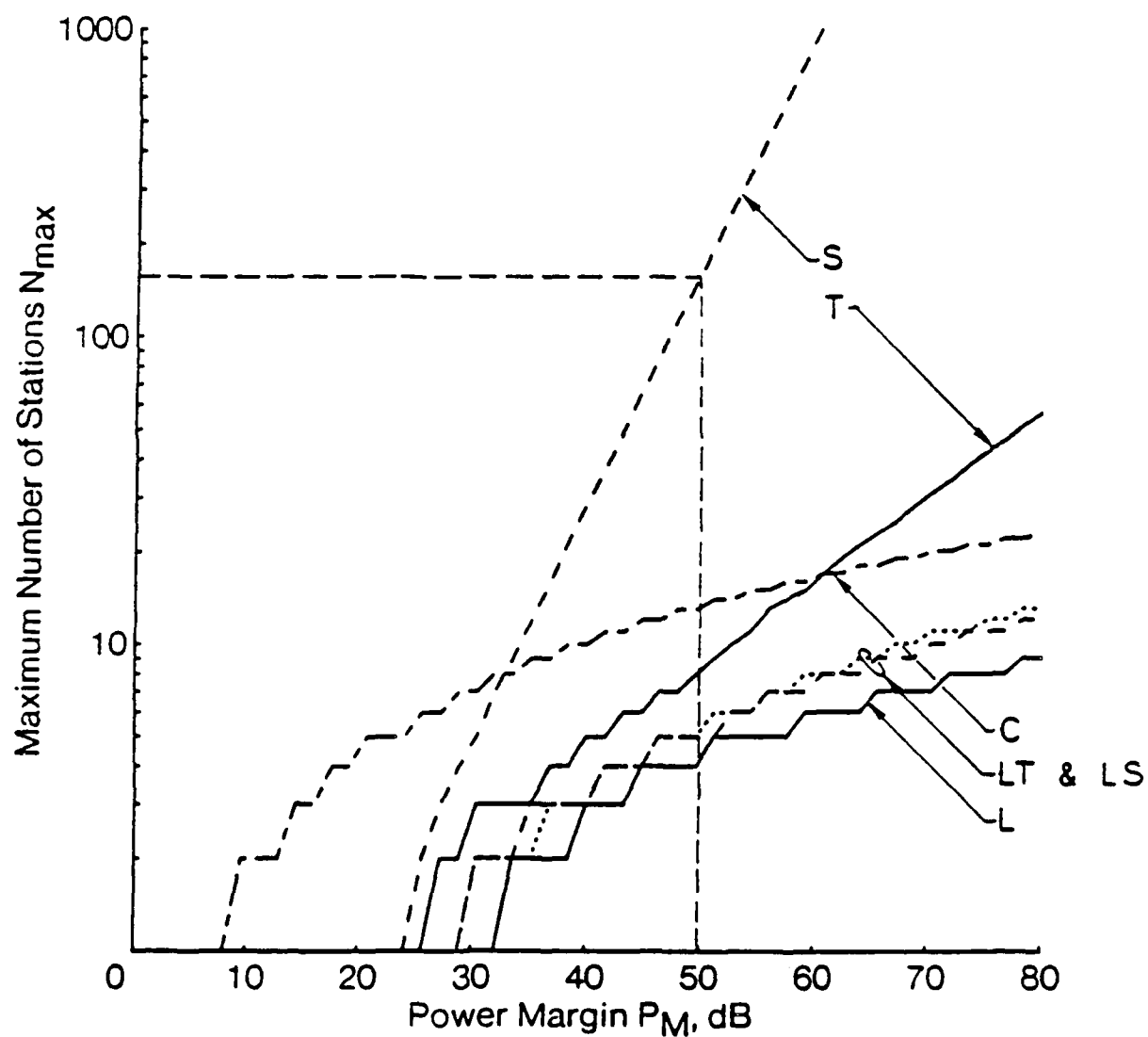


Fig. A.20 The maximum number of stations as a function of the power margin for LC components, multimode fiber, and individually optimized couplers.

Figures A.18, A.19, and A.20 display $N_{\max}$ as a function of the power margin P_M for HC, MR, and LC component parameters (except for P_M), respectively. In these figures we have assumed multimode fiber and individually optimized couplers. We note that in T and S subnetworks, the rate of increase in $N_{\max}$ as a function of P_M is about one decade per 20 and 10 dB, respectively. This result is consistent with the bounds (4.16) and (4.17) obtained in chapter 4, section 4.3.3. In the other subnetworks, however, this rate of increase is below the ideal rate of one decade per 20 dB, due to the fact that the number of couplers between transmitters and receivers grows linearly with N . From these figures we also note the importance of P_M . When P_M is larger than 60 dB, and all other parameters have their MR values, all configurations can support more than 20 stations. In fact for a power margin of just 50 dB, when all other parameters assume their LC values, the S subnetwork can still support about 150 stations. (The figures can also be used to compute $N_{\max}$ of a C subnetwork when we take the difference between receiver and sensor sensitivities into account.)

For the uniform optimization the behavior is the same as above. Fig. A.21 displays $N_{\max}$ as a function of margin P_M for HC components, multimode fiber, and uniformly optimized couplers. In the case of T subnetwork, $N_{\max}$ is almost always equal to some power of two (hence the shape of the curve) for the following reason:* for simplicity consider ideal components; when the coupling fraction is $\frac{1}{2}$, the minimum transmittance from transmitters to the funneling point (or from the funneling point to receivers) is $(\frac{1}{2})^{\lceil \log_2 N \rceil}$; or the minimum transmittance between transmitters and receivers is $(\frac{1}{2})^{2\lceil \log_2 N \rceil}$; therefore $N_{\max}$ is equal to some power of 2. In the case of LT subnetwork, due to the linear collection subnetwork the transitions between the successive powers of two are more gradual than those in T

*To show more clearly the transitions of $N_{\max}$ between the powers of two, we have generated 1000 data points for the T and LT curves in Fig. A.21. All other curves contain 50 data points.

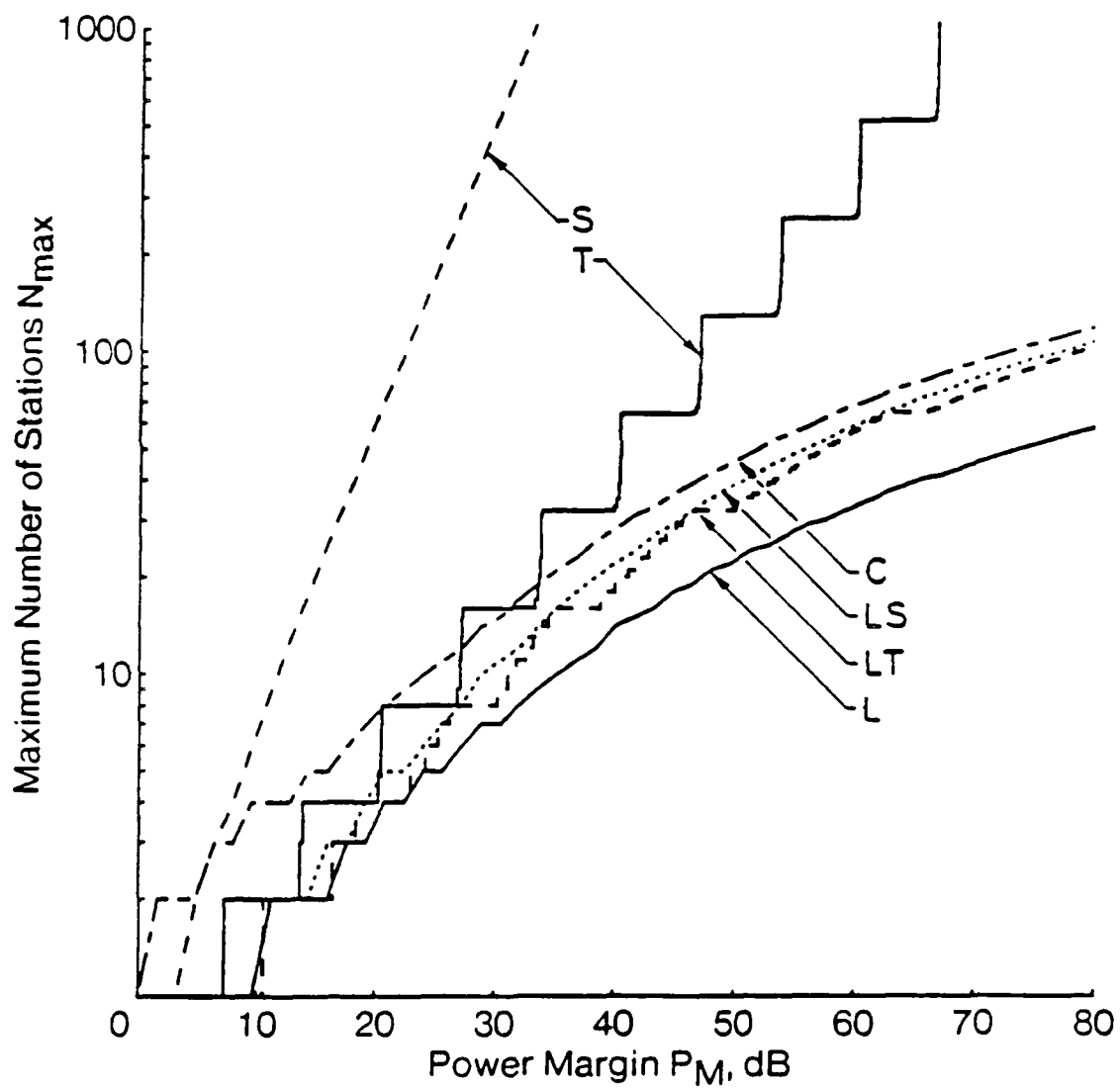


Fig. A.21 The maximum number of stations as a function of the power margin for HC components, multimode fiber, and uniformly optimized couplers.

subnetwork.

Fig. A.22 displays $N_{\max}$ as a function of the coupler excess-loss E_C . In this figure (and subsequent figures), we have assumed multimode fiber, individually optimized couplers, MR components in the S subnetwork, and HC components in the other subnetworks. for multimode fiber and individually optimized components. We observe that, $N_{\max}$ decreases significantly as E_C increases. Note that in the T and S subnetworks the rate of decrease is constant over the entire domain of $N_{\max}$, while in the subnetworks with a linear collection subnetwork, the rate of decrease is larger at smaller larger $N_{\max}$. The reason is that, in the T and S subnetworks, the average number of couplers (and hence their excess-loss) on the transmitter-to-receiver paths is on the order of $\log_2 N$, while in the subnetworks with a linear collection subnetwork it is on the order of N .

In all subnetworks other than S, $N_{\max}$ decreases considerably as the joint insertion loss L_J increases (see Fig. A.23); in the S subnetwork, regardless of the number of stations, there are only a few joints on the path from every transmitter to every receiver. As shown in Fig. A.24, when the fiber attenuation f is increased, only the S subnetwork suffers a significant decrease in $N_{\max}$; in the other subnetworks, the losses due to fiber attenuation are negligible compared to the other losses.

$N_{\max}$ is not sensitive to L_C (Based on the convention we have used here, L_C only denotes the insertion losses of the connectors at transmit and receive terminals. The insertion losses of the other connectors in the network are denoted by L_J .) For the values of M_C considered here ($< 5\%$), $N_{\max}$ is not sensitive to M_C .

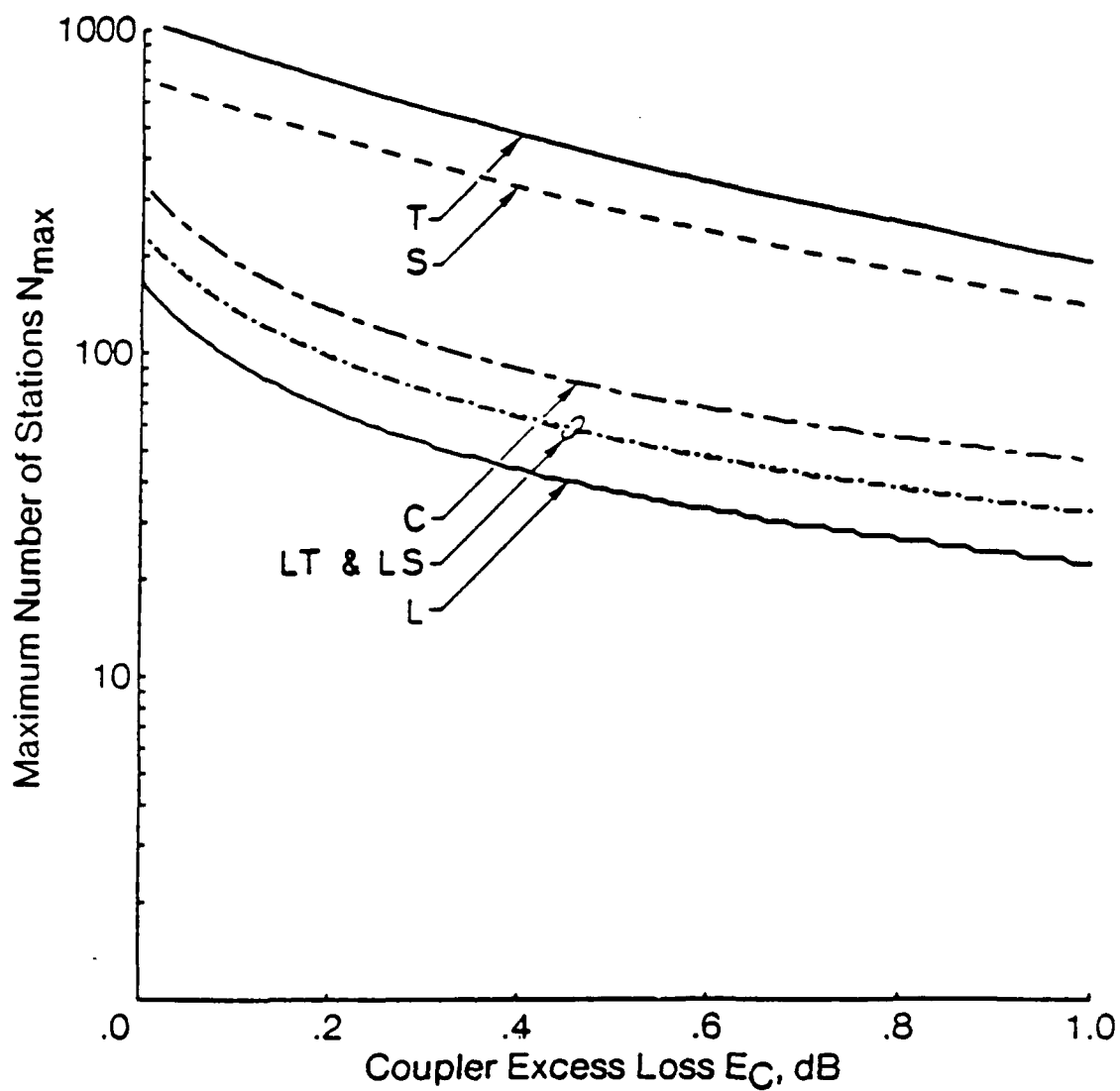


Fig. A.22 The maximum number of stations as a function of the coupler excess-loss for multimode fiber and individually optimized couplers. In the S subnetwork MR components are assumed while in the other subnetworks HC components are assumed.

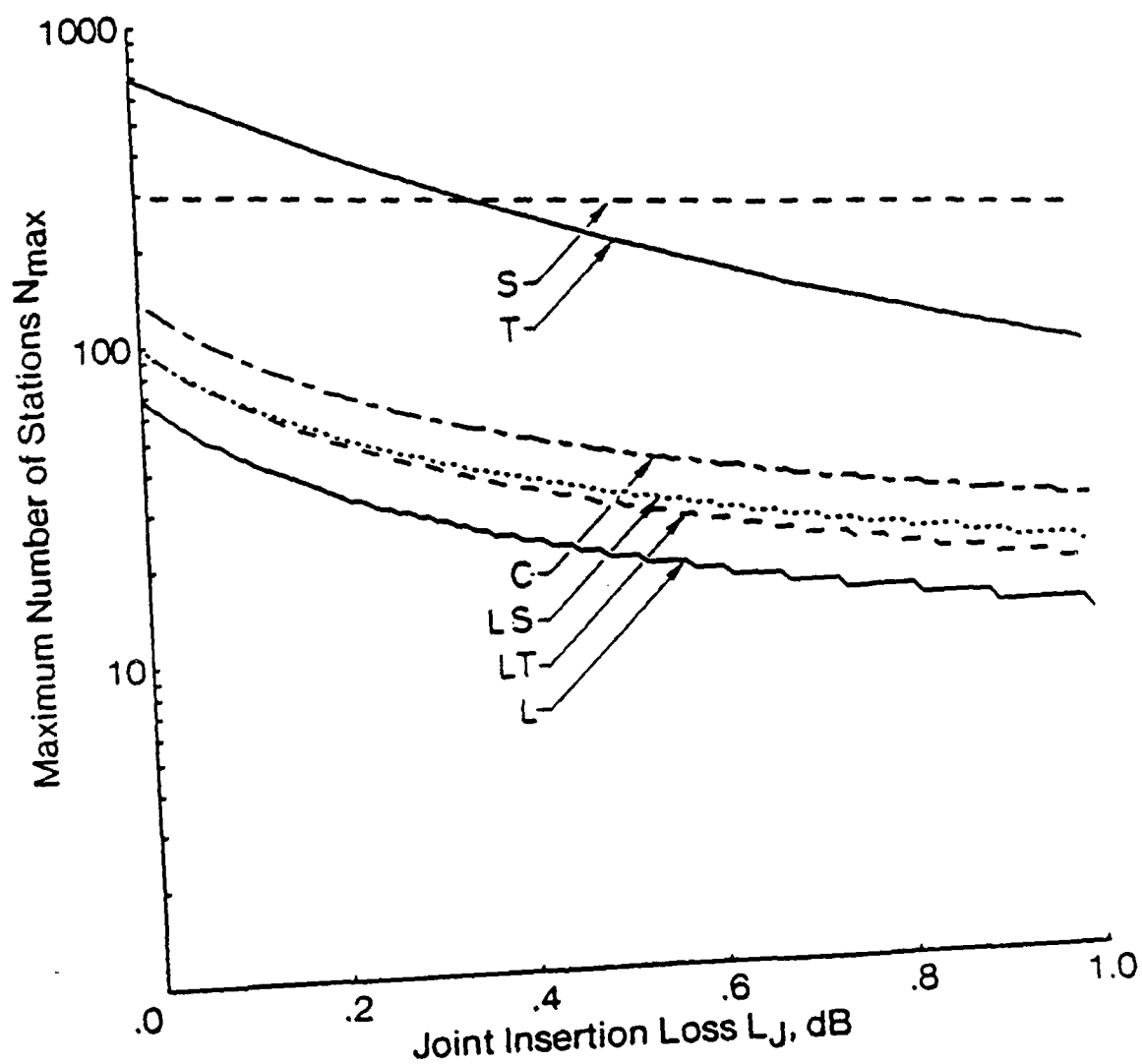


Fig. A.23 The maximum number of stations as a function of the joint insertion loss for multimode fiber and individually optimized couplers. In the S subnetwork MR components are assumed while in the other subnetworks HC components are assumed.

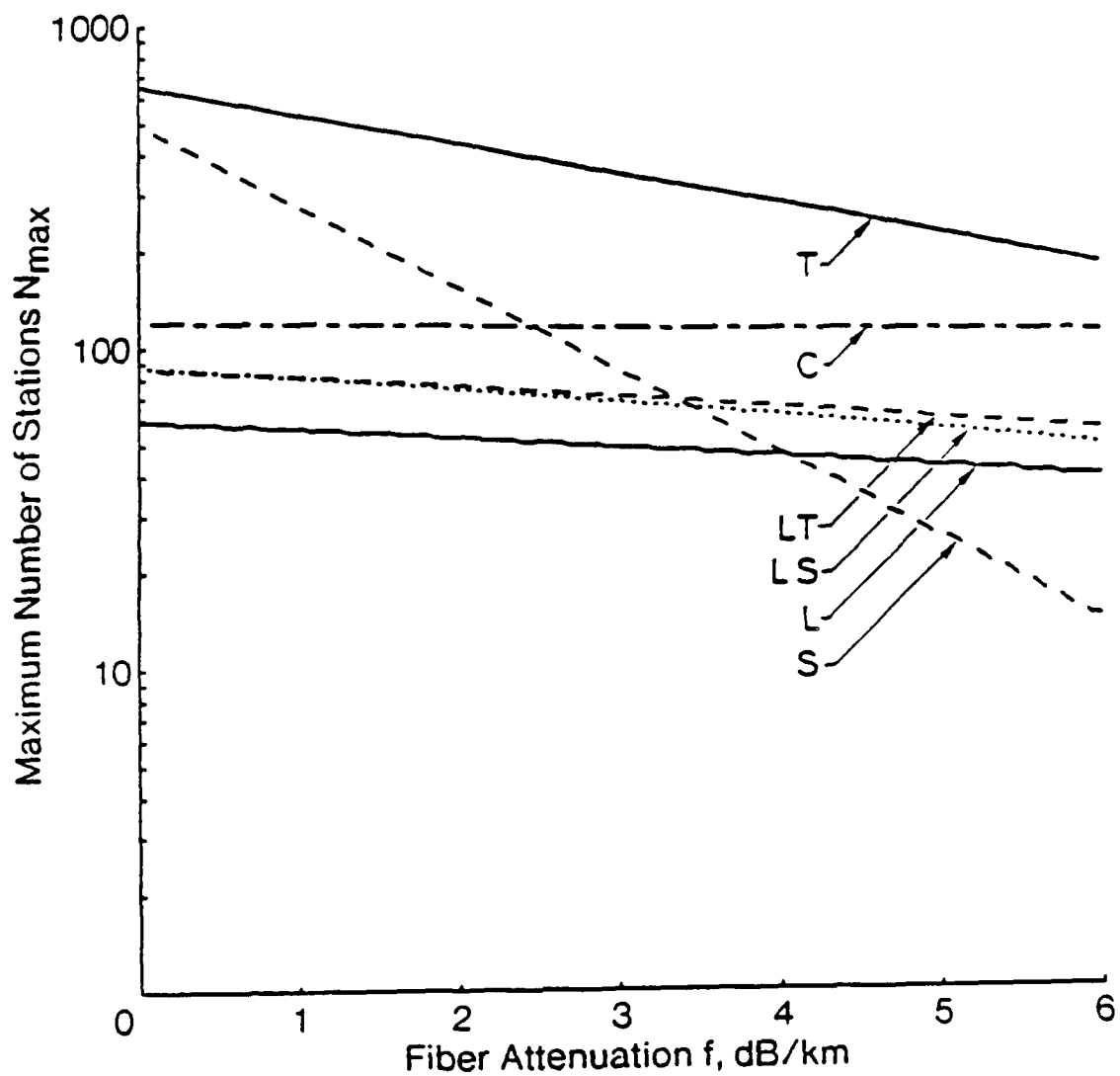


Fig. A.24 The maximum number of stations as a function of the fiber attenuation for multimode fiber and individually optimized couplers. In the S sub-network MR components are assumed while in the other subnetworks HC components are assumed.

References

- [ALTM77] D. E. Altman and H. F. Taylor, "An eight-terminal fiber optics data bus using tee couplers," *Fiber and Integrated Optics*, Vol. 1, 1977, pp. 135-152.
- [AURA77] F. Auracher and H.-H. Witte, "Optimized layout for a data bus system based on a new planar access coupler," *Applied Optics*, Vol. 16, No. 12, December 1977, pp. 3140-3142.
- [BARA64] P. Baran, "On distributed communications: XI. Summary overview." Rand Corporation, Santa Monica, California, Memorandum RM-3767-PR, August 1964.
- [BUX] W. Bux, "Local-Area Subnetworks: A Performance Comparison." *IEEE Transactions on Communications*, Vol. COM-29, No. 10, October 1981, pp. 1465-1473.
- [CANSTAR] Canstar Communications, Model PCS.
- [CHLA83] I. Chlamtac and M. Eisinger, "Voice/Data integration on Ethernet: Backoff and priority considerations," Technical Report 273. Department of Computer Science, Technion, Israel Institute of Technology, Haifa, Israel, May 1983.

- [CHLA85] I. Chlamtac and M. Eisinger, "Performance of Integrated Services (voice/data) CSMA/CD networks," in *ACM Sigmetrics Conference on Measurement and Modeling of Computer Systems*. Austin, Texas. August 1985, pp. 87-93.
- [CLAR78] D. D. Clark, K. T. Pograd, D. P. Reed, "An introduction to local area networks," *Proceedings of the IEEE*, Vol. 66, No. 11, November 1978.
- [CRAN75] M. A. Crane and D. L. Iglehart, "Simulating stable stochastic systems: III. Regenerative processes and discrete event simulation." *Operations Research*, 23, pp. 33-45.
- [DAVI67] D. W. Davies, K. A. Bartlett, R. A. Scantlebury, and P. T. Wilkinson, "A digital communication network for computers giving rapid response at remote terminals." *ACM Symposium on Operating System Principles*, Gatlinburg, Tennessee, October 1967.
- [DETR84] J. D. DeTreville, "A simulation-based comparison of voice transmission on CSMA/CD networks and on token buses," *AT&T Bell Laboratories Technical Journal*, Vol. 63, No. 1, January 1984, pp. 33-55.
- [FINE84] M. Fine and F. A. Tobagi, "Demand Assignment Multiple Access Schemes in Broadcast Bus Local Area Networks," *IEEE Transactions on Computers*, Vol. C-33, No. 12, December 1984, pp. 1130-1159.
- [FINE85a] M. Fine, "Performance of demand assignment multiple access schemes in broadcast bus networks," Ph.D. Thesis, Department of Electrical Engineering, Stanford University, Stanford, California, June 1985.
- [FINE85b] M. Fine and F. A. Tobagi, "Packet voice on a local area network with round robin service," *Technical Report SEL 85-275*, Computer Systems Laboratory, Stanford University, Stanford, California, April

1985.

- [FRAT81] L. Fratta, F. Borgonovo and F. A. Tobagi. "The Express-Net: A Local Area Communication Network Integrating Voice and Data." in *Performance of Data Communication Systems*, G. Pujolle (editor). North Holland, Amsterdam, 1981, pp. 77-88.
- [FREN82] S. French, *Sequencing and Scheduling: An Introduction to the Mathematics of the Job-Shop*. John Wiley, New York, 1982.
- [GALL78] R. G. Gallager, "Conflict resolution in random access broadcast networks," presented at *AFOSR Workshop on Communication Theory*. Provincetown, Massachusetts, 1978.
- [GERL83a] M. Gerla, C. Yeh and P. Rodrigues. "A token protocol for high speed fiber optics local networks." in *Proceedings of Conference on Optical Fiber Communication*, New Orleans, Louisiana, February 1983.
- [GERL83b] M. Gerla, P. Rodrigues and C. Yeh, "BUZZ-NET: A hybrid random access/virtual token local network," in *Proceedings of Globecom '83*. San Diego, CA, December 1983, pp. 4.35.1-5.
- [GONS83] T. A. Gonsalves, "Packet-voice communication on an Ethernet local network: An experimental study," in *ACM SIGCOM Symposium on Communications Architectures and Protocols*, Austin, Texas, March 1983, pp. 178-185.
- [HUDS74] M. C. Hudson and F. L. Thiel, "The Star Coupler: A unique interconnection component for multimode optical waveguide communications systems," *Applied Optics*, Vol. 13, No. 11, November 1974, pp. 2540-2545.
- [HANS81] K. Hanson, W. Chou and A. Nillson. "Integration of Voice, Data,

- and Image Traffic on a Wideband Local Network." in *Proceedings of Computer Network Symposium*, December 1981.
- [IEEE83] *IEEE Project 802 Local Area Network Standards*. Draft D 802.4. Token-Passing Bus Access Method and Physical Layer Specification. IEEE Computer Society, Silver Spring, MD, 1983.
- [IEEE85a] "Carrier Sense Multiple Access Method and Physical Layer Specifications." *ANSI/IEEE Standard 802.3*. New York, New York, 1985.
- [IADA80] I. Iiada, M. Ishizuka, Y. Yasuda and M. Onoe, "Random access packet switched local computer network with priority function." in *NTC80*. Houston, Texas, December 1980, pp. 37.4.1-37.4.6.
- [JACO78] I. M. Jacobs, R. Binder and E. V. Hoversten, "General purpose packet satellite networks." in *Proceedings of IEEE*. Vol. 66, No. 11, November 1978, pp. 1448-1467.
- [JOHN81] D. H. Johnson and G. C. O'Learly, "A local access network for packetized digital voice communication," *IEEE Transactions on Communications*, Vol. COM-29, No. 5, May 1981, pp. 679-688.
- [KLEI75] L. Kleinrock 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, December 1975.
- [KLEI76] L. Kleinrock, *Queueing Systems, Volume I: Theory*, John Wiley & Sons, 1976.
- [KLIM74] G. P. Klimov, "Time-sharing service systems. I, Theory of probability and applications," Vol. 19, No. 3, 1974, pp. 532-551.
- [KLIM78] G. P. Klimov, "Time-sharing service systems. II. Theory of probabil-

ity and applications." 1978, pp. 314-21.

- [KUMMS2] K. Kummerle and M. Reiser, "Local area communication networks - an overview," *Journal of Telecommunication Networks*, Vol. 1, No. 4, Winter 1982.
- [KURO82] J. F. Kurose and M. Schwartz, "A family of window protocols for time constrained applications in CSMA networks," *INFOCOM'83*, San Diego, California, April 1983.
- [LI78] T. Li, "Optical fiber communication - The state of the art," *IEEE Transactions on Communications*, Vol COM-26, July 1978, pp. 946-955.
- [LIMB82] J. O. Limb and C. Flores, "Description of Fasnet, a unidirectional local area communications network," *Bell Systems Technical Journal*, Vol. 61, No. 7, September 1982, pp. 1413-1440.
- [LIMB83] J. O. Limb and L. E. Flamm, "A distributed local area network protocol for combined voice and data transmission," *IEEE Journal on Selected Areas in Communications*, Vol. SAC-1, No. 5, November 1983, pp. 926-934.
- [LIMB84] J. O. Limb, "On Fiber Optic Taps for Local Area Networks," in *Proceedings of International Conference on Communications*, Amsterdam, May 1984, pp. 1130-1136.
- [MARH84] M. E. Marhic, "Combinatorial star couplers for singlemode optical fibers," *Optics Letters*, Vol.9, No.8, August 1984, pp. 368-370.
- [MARS82] M. A. Marsan and G. Albertengo, "Integrated voice and data network," *Computer Communications*, Vol. 5, No. 3, June 1982.
- [MAXE82] N. F. Maxemchuk, "A variation on CSMA/CD that yields movable

TDM slots in integrated voice/data local networks." *Bell System Technical Journal*, Vol. 61, No. 7, September 1982, pp. 1527-1550.

- [METC76] R. M. Metcalfe and D. R. Boggs, "Ethernet: Distributed packet switching for local computer networks," *Communications of the ACM*, Vol. 19, No. 7, July 1976, pp. 395-403.
- [MIDW82] J. E. Midwinter, "First-generation trunk transmission Systems: Capabilities and limitations," *IEEE Journal on Selected Areas in Communications*, Vol. SAC-1, April 1983, pp. 381-386.
- [MILT76] A. F. Milton and A. B. Lee, "Optical access couplers and a comparison of multiterminal fiber communication systems," *Applied Optics*, Vol. 15, No. 1, January 1976, pp. 244-252.
- [MUOI84] T. V. Muoi, "Receiver Design for High-Speed Optical-Fiber Systems," *Journal of Lightwave Technology*, Vol. LT-2, No. 3, June 1984, pp. 243-267.
- [NUTT82] G. J. Nutt and D. L. Bayer, "Performance of CSMA/CD networks under combined voice and data loads," *IEEE Transactions on Communications*, Vol. COM-30, No. 1, January 1982, pp. 6-11.
- [OGAW83] K. Ogawa, "Considerations for optical receiver design," *IEEE Journal on Selected Areas in Communications*, Vol. SAC-1, No. 3, April 1983, pp. 533-540.
- [RHOD83] N. L. Rhodes, "Interaction of network design and fiber optic component design in local area networks," *IEEE Journal on Selected Areas in Communications*, Vol. SAC-1, No. 3, April 1983, pp. 489-492.
- [RIVE84] P. Rivett, P. Eng. and K. Eastwood, "Asymmetric Couplers for Linear Networks: Part 1," in *Fiber Optics Now*, Vol. 6, No. 3, Winter 1984.

- [RODR84] P. Rodrigues, L. Fratta and M. Gerla, "Token-less protocols for fiber optics local area networks," in *Proceedings of International Conference on Communications*, Amsterdam, May 1984, pp. 1150-1153.
- [SCHM83] R. V. Schmidt, E. G. Rawson, R. E. Norton, S. B. Jackson and M. D. Douglas, "Fibernet II: A Fiber Optic Ethernet", *IEEE Journal on Selected Areas in Communications, Special Issue on Local Area Networks*, Vol. SAC-1, No. 5, November 1983, pp. 702-711.
- [STAL84] W. Stallings, *Local Area Networks: An Introduction*, New York: Macmillan, 1984.
- [TANE81] A. S. Tanenbaum, *Computer Networks*, Prentice Hall, 1981.
- [TOBA79] F. Tobagi and V. B. Hunt, "Performance analysis of carrier sense multiple access with collision detection," in *Proceedings of the Local Area Communications Network Symposium*, Boston, Massachusetts, May 1979.
- [TOBA80] F. A. Tobagi, "Multiaccess protocols in packet communication systems," *IEEE Transactions on Communications*, Vol. COM-28, No. 4, April 1980.
- [TOBA82a] F. A. Tobagi, "Carrier sense multiple access with message-based priority functions," *IEEE Transactions on Communications*, Vol. COM-30, No. 1, January 1982, pp. 185-200.
- [TOBA82b] F. A. Tobagi and N. Gonzalez-Cawley, "On CSMA-CD local networks and voice communication," in *INFOCOM'82*, Las Vegas, Nevada, March/April 1982.
- [TOBA83] F. A. Tobagi, F. Borgonovo and L. Fratta, "Express-net: A high-performance integrated-services local area network," *IEEE Journal*

on Selected Areas in Communications, Vol. SAC-1, No. 5, November 1983, pp. 898-912.

- [TOWS82] D. Towsley and G. Venkatesh. "Window random access protocols for local computer networks." *IEEE Transactions on Computers*, Vol. C-31, No. 8, August 1982, pp. 715-722.
- [VILL81] C. A. Villarruel, R. P. Moeller and W. D. Burns. "Tapped tee single-mode data distribution system." *IEEE Journal of Quantum Electronics*, Vol. QE-17, No. 6, June 1981, pp. 941-946.
- [WOOD85] T. H. Wood and M. S. Whalen. "Demonstration of effectively non-reciprocal optical fiber directional couplers," in *Proceedings of Conference on Optical fiber Communication*, San Diego, CA, February 1985, pp. PD14.1-4.