

INFORMATION TO USERS

This manuscript has been reproduced from the microfilm master. UMI films the text directly from the original or copy submitted. Thus, some thesis and dissertation copies are in typewriter face, while others may be from any type of computer printer.

The quality of this reproduction is dependent upon the quality of the copy submitted. Broken or indistinct print, colored or poor quality illustrations and photographs, print bleedthrough, substandard margins, and improper alignment can adversely affect reproduction.

In the unlikely event that the author did not send UMI a complete manuscript and there are missing pages, these will be noted. Also, if unauthorized copyright material had to be removed, a note will indicate the deletion.

Oversize materials (e.g., maps, drawings, charts) are reproduced by sectioning the original, beginning at the upper left-hand corner and continuing from left to right in equal sections with small overlaps. Each original is also photographed in one exposure and is included in reduced form at the back of the book.

Photographs included in the original manuscript have been reproduced xerographically in this copy. Higher quality 6" x 9" black and white photographic prints are available for any photographs or illustrations appearing in this copy for an additional charge. Contact UMI directly to order.

UMI

**A Bell & Howell Information Company
300 North Zeeb Road, Ann Arbor MI 48106-1346 USA
313/761-4700 800/521-0600**

All-Optical Networks with Sparse Wavelength Conversion

by

Suresh Subramaniam

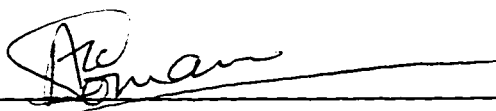
A dissertation submitted in partial fulfillment of
the requirements for the degree of

Doctor of Philosophy

University of Washington

1997

Approved by



(Chairperson of Supervisory Committee)

Program Authorized

to Offer Degree

Electrical Engineering

Date

8/7/97

UMI Number: 9807034

UMI Microform 9807034
Copyright 1997, by UMI Company. All rights reserved.

**This microform edition is protected against unauthorized
copying under Title 17, United States Code.**

UMI
300 North Zeeb Road
Ann Arbor, MI 48103

In presenting this thesis in partial fulfillment of the requirements for a doctoral degree at the University of Washington, I agree that the library shall make its copies freely available for inspection allowable only for scholarly purposes, consistent with "fair use" as prescribed in the U.S. Copyright Law. Requests for copying or reproduction of this dissertation may be referred to University Microfilms, 1490 Eisenhower Place, P. O. Box 975, Ann Arbor, MI 48106, to whom the author has granted "the right to reproduce and sell (a) copies of the manuscript in microform and/or (b) printed copies of the manuscript made from the microform."

Signature 

Date 8/14/'97

University of Washington

Abstract

All-Optical Networks with Sparse Wavelength Conversion

by Suresh Subramaniam

Chairperson of Supervisory Committee: *Professor Arun Somani*

Electrical Engineering

Wavelength conversion is considered a key capability for improving the efficiency of wavelength-routed, wavelength-division-multiplexed (WDM) all-optical networks. We consider networks with an arbitrary number of wavelength converters and develop accurate analytical models of reasonable complexity for evaluating the blocking performance under a variety of dynamic input traffic. We also consider the problem of optimally placing a given number of wavelength converters on a path of a network and develop a dynamic programming algorithm to solve it. This framework is used to obtain the optimal converter placement in bus and ring topologies. Finally, we consider the problem of wavelength assignment in fixed-routing networks with no wavelength conversion and develop efficient heuristics.

TABLE OF CONTENTS

List of Figures	iii
Chapter 1: Introduction	1
Chapter 2: Wavelength routing	5
2.1 Wavelength Routing	5
2.2 Routing Node Architectures	6
2.3 Wavelength Conversion	8
2.4 Contributions and Outline of the Dissertation	12
Chapter 3: An analytical model for sparse wavelength conversion: Poisson traffic	15
3.1 An Analytical Model for Path Blocking Performance	17
3.2 Sparse Wavelength Conversion	26
3.3 Blocking in a Ring Network	31
3.4 Blocking in a Mesh-Torus Network	34
3.5 Blocking in a Hypercube Network	39
3.6 Blocking in Random Topologies	42
3.7 Conclusions	47
Chapter 4: An analytical model for sparse wavelength conversion: Non-Poisson traffic	49
4.1 Introduction	49

4.2	The Single Link Model	50
4.3	Network Blocking Analysis	57
4.4	Performance Results	60
4.5	Conclusions	65
Chapter 5:	On the optimal placement of wavelength converters	67
5.1	Introduction	67
5.2	Converter Placement in a Path	70
5.3	Converter Placement in a Bus	86
5.4	Converter Placement in a Ring	91
5.5	The Effect of Traffic Model on Converter Placement	96
5.6	Conclusions	98
Chapter 6:	Dynamic wavelength assignment in fixed-routing net-	
	works	100
6.1	Introduction	100
6.2	The Problem and the $M\Sigma$ Algorithm	102
6.3	Implementing the $M\Sigma$ Algorithm	104
6.4	An Analytical Blocking Model for Dense Multi-fiber Networks	107
6.5	Results	109
6.6	Conclusions	114
Chapter 7:	Conclusions and future work	115
Bibliography		119
Appendix A:	An ILP for converter placement	126

LIST OF FIGURES

2.1	A wavelength-routing network.	6
2.2	A passive wavelength-routing node.	7
2.3	A configurable wavelength-routing node.	8
2.4	A wavelength-routing network with wavelength conversion.	9
2.5	A routing node with full wavelength conversion.	10
2.6	The share-per-node wavelength-convertible switch architecture.	11
2.7	The share-per-link wavelength-convertible switch architecture.	11
3.1	Calls arriving and leaving on a two-hop path.	20
3.2	The state-space of the three-dimensional Markov chain.	21
3.3	An l -hop path with nodes numbered $s, 1, 2, \dots, l-1, d$	27
3.4	The blocking probability vs. the load per station for a 100-node ring network. (a) $F = 5$, and (b) $F = 20$	33
3.5	The blocking probability vs. the converter density for a ring network with a total network load of 2 Erlangs. Solid lines: $N = 20$ and $\rho = 0.1$, Dashed lines: $N = 100$ and $\rho = 0.02$	34
3.6	The converter density vs. the number of stations that can be supported for a ring network. $F = 5$, $P_b = 10^{-3}$	35
3.7	A 5×5 bidirectional mesh-torus network.	36
3.8	The blocking probability vs. the load per station for a 101×101 mesh-torus network with 5 wavelengths per fiber.	37

3.9	The blocking probability vs. conversion density for (a) a 11×11 mesh-torus with $\rho = 0.5$, and (b) a 101×101 mesh-torus with $\rho = 0.01$	38
3.10	The converter density vs. the square root of the number of stations that can be supported for a bidirectional mesh-torus network. $F = 5, P_b = 10^{-3}$	39
3.11	The blocking probability vs. the load per station for a hypercube network when the number of wavelengths per fiber is 5. (a) $N = 32$, and (b) $N = 1024$	40
3.12	The blocking probability vs. the converter density for a 1024-node hypercube network with a load of 0.1 Erlangs per station.	41
3.13	The blocking probability vs. the converter density for a random network. $N = 100, b = 10, \rho = 1$	44
3.14	The blocking probability vs. the average out-degree per node for a random network. $N = 100, F = 4, \rho = 1$	45
3.15	The blocking probability vs. the load per station for a random network. $N = 100, F = 4, b = 10$	45
3.16	The number of stations supported vs. the converter density for a random network. $F = 4, b = 10$	46
3.17	The blocking probability vs. the number of wavelengths per link for a random network. $N = 500, F \times b = 64$	46
4.1	The Markov chain for the Pascal model.	53
4.2	The Markov chain for the Bernoulli model.	54
4.3	The blocking probability vs. the load per station for the mesh network with no conversion and full conversion. (a) $Z = 1.0$ (b) $Z = 1.25$ (c) $Z = 1.5$	63
4.4	The blocking probability vs. the conversion density with a load of 1.0 per station for (a) the mesh, and (b) the hypercube.	64
4.5	The blocking probability vs. the link peakedness for the hypercube. $\rho_s = 1$	66

4.6	The conversion gain vs. the link peakedness for the mesh network. $P_b = 10^{-3}$	66
5.1	A path of length H	70
5.2	P_b vs. ρ for optimal/uniform (O) and random (R) placement with $K = 1$ and 2 converters.	74
5.3	P_b vs. K for optimal/uniform (O), random (R_r) placement of K converters, and random placement of a random number of converters (R_q) with mean K , for uniform link loads ($\rho = 0.1$).	76
5.4	G_u^r and G_u^q vs. K for $P_b \approx 10^{-3}$	77
5.5	The bound for G_u^q vs. F for $P_b = 10^{-3}$ and $K = \lceil 0.2(H - 1) \rceil$	79
5.6	P_b vs. the converter position for $K = 1$ for uniform link loads ($\rho = 0.1$) and for linearly increasing link loads ($\rho_0 = 0.05, \rho_9 = 0.1$).	82
5.7	P_b vs. K for random (R), uniform (U) and optimal (O) placements. Link loads are linearly increasing with $\rho_0 = 0.125, \rho_9 = 0.25$	83
5.8	P_b vs. K for random (R), and optimal (O) placements. $H = 10, F = 10$. $\rho_i = 0.01(H - i)(i + 1)$. Solid lines: End-to-end performance, Dashed lines: Average performance.	86
5.9	A bus of length H	88
5.10	P_b vs. K for uniform link loads, $\rho = 0.4$. $H = 10, F = 10$	90
5.11	P_b vs. converter position for $K = 1$ and uniform link loads, $\rho = 0.3$. $H = 10, F = 10$	90
5.12	Traffic vs. node number for uniform link loads, $\rho = 0.3$	91
5.13	(a) The ring logically broken at node i where a converter is assumed to be placed. (b) The resulting bus after relabeling (corresponding old labels are shown below the nodes).	92
5.14	P_b vs. K for uniform link loads, $\rho = 0.4$. $H = 10, F = 10$	95

5.15	P_b of an end-to-end call vs. K for a 10-hop path with optimal/uniform (O), random (R_r) placement of K converters, and random placement of a random number of converters (R_q) with mean K . Poisson traffic model.	97
5.16	P_b vs. the converter position for $K = 1$ for uniform link loads ($\rho = 0.1$) and for linearly increasing link loads ($\rho_0 = 0.05, \rho_9 = 0.1$). Poisson traffic model.	98
6.1	P_b vs. F for the different algorithms in a 20-node unidirectional ring (a) $d = 1$, load/f/d = 1 Erlang (b) $d = 10$, load/f/d = 1.6 Erlangs.	110
6.2	P_b vs. F for the different algorithms in a 5×5 bidirectional mesh-torus (a) $d = 1$, load/f/d = 25 Erlangs (b) $d = 3$, load/f/d = 31 Erlangs.	111
6.3	L_p vs. d for a 5×5 mesh-torus. $P_b(WC) \approx 0.0001$, $F = 10$	112
6.4	L_p vs. d for a 20-node ring. $P_b(WC) \approx 0.0001$ (a) $F = 10$ (b) $F = 60$	113

ACKNOWLEDGMENTS

This is quite possibly the hardest part of this document for me to write. It is impossible to thank everyone responsible for the completion of this dissertation enough.

I am thankful to my advisor Professor Arun K. Somani for helping me in several ways during the last three years. He has always given me his valuable time and advice whenever I needed them from him. I also appreciate the freedom he allowed me in pursuing my research ideas.

Professor Murat Azizoglu has played a large role in shaping this dissertation. I am grateful to him for the utmost care and patience with which he has read every technical document I wrote in the last three years. His feedback and suggestions have significantly helped improve this thesis.

I would also like to thank the other members of my committee, Professors Eve A. Riskin, Richard E. Ladner, and James S. Meditch, for their comments on a draft of this thesis.

The work reported in Chapter 6 was performed at MIT Lincoln Laboratory during the Summer of 1996, jointly with Dr. Richard Barry. I thank him for giving me the opportunity to work in his group.

I acknowledge the financial support provided by NSF under grants 9103485 and 9502610, and the University of Washington Royalty Research Fund.

This thesis would not have materialized without the help and support of friends and family. I have been fortunate to have a number of good friends in

the department, and an exhaustive list would be too long. However, I must mention Sathyadev Uppala and Govind Krishnamurthi for the many good laughs we have shared. I would also like to thank Anand Bedekar with whom I have discussed over tea, topics ranging from queueing theory to basketball to Indian politics. Thanks are also due to Siva and his family, and Prakash and his family, for the wonderful time we had together at Sand Point while having dinner and playing cards.

My dear wife Deepa has been a model of patience and understanding over the last three years. There have been many days on which I was immersed in thought and hardly spoke to her. Yet, at the end of the day, she would cheerfully ask what work I had completed that day and would only urge me to work harder. Her encouragement has been a constant source of inspiration to me.

My parents feel pride in every little thing I achieve. Though they may never admit it, I am sure they are very proud that all three of their sons are now doctoral degree holders. My adoptive parents, who are no longer in this world, would have felt equally proud. My brothers and sister will undoubtedly be very happy for me.

My son Ashwin was born on July 26, five days before my thesis defense, and made that week my most memorable to date. This thesis is dedicated to my parents and the little one.

Chapter 1

INTRODUCTION

The evolution of bandwidth demand per user has taken many forms over the years. To consider an example, the per user bandwidth demand for World Wide Web access alone has increased eight-fold annually over the last few years [1]. Furthermore, applications such as medical image transfer, supercomputer visualization, and multimedia teleconferencing require throughputs on the order of a few Gb/s per user and a few Tb/s per network. It is projected that the number of gigabit applications will rise sharply in the future [1]. It has long been recognized that optical fiber with its almost limitless bandwidth (about 25 THz or 200 nm in the $1.5\mu m$ band for single-mode optical fiber) is the most suitable transmission medium for such information transfer rates. Other characteristics of optical fiber such as the provision of almost complete immunity from noise (bit error rates of 10^{-15} are achievable) and tight packing density further enhance its attractiveness. Anticipating the potential benefits of optical fiber, millions of miles of optical fiber have been installed by long-haul transmission companies and regional telephone companies. With the development of important optical access technologies, much of it is being used as point-to-point transmission links today. These links typically operate at data rates of a few Mb/s and use simple on-off keyed (OOK) intensity modulation of lasers and direct detection receivers [2].

Likewise, in computer networks, the current generation of optical networks use optical fiber instead of conventional copper wires as the transmission medium (e.g., Distributed Queue Dual Bus (DQDB) and Fiber Distributed Data Interface (FDDI)

which have network throughputs ≈ 100 Mb/s). These networks are also called electro-optic networks [3] or multi-hop networks [4]. In these networks, information is transmitted from one node to another via optical fiber, converted to electronic form and processed, converted back to optical form, and transmitted to another node, until the destination node is reached. Even though the use of optical fiber has enabled higher transmission speeds and lower error rates than achievable using copper wires, the speed of transmission is limited by the speed of electronics at the network nodes. Because the total capacity in such networks is time-shared among many users, each user has to have electronics that run at the aggregate speed of the network, rather than at the speed of a user. However, the speed of electronic digital circuitry (≈ 5 Gb/s for GaAs) is such that Tb/s networks cannot be realized. This constraint due to electronics is the infamous “electronic bottleneck” [3].

All-optical networks (AONs) get around this bottleneck by limiting optoelectronic conversion to the periphery of the network (the user nodes). While commercial optical technology is mature enough for the deployment of optical links, there are several technological limitations that make an all-optical network conceptually very different from an electronic network. For instance, the lack of optical memory and the ability of optical technology to perform only primitive routing functions means that packet switching as we know it cannot be performed by optical networks today. The challenge in optical networking over the years has been to develop clever architectural solutions that hide the technology limitations while achieving the high bandwidths that applications demand.

The bandwidth of optical fiber far exceeds any single user’s requirements. Several multiple access schemes are available for partitioning the fiber bandwidth among many users. Time Division Multiple Access (TDMA) and Code Division Multiple Access (CDMA) require very high speed optical and/or electronic circuitry for transmission and synchronization. Currently, the most popular way to effectively access the fiber bandwidth is Wavelength Division Multiplexing (WDM) [3] in which the

available optical spectrum is partitioned into several channels, each corresponding to a different wavelength (frequency). In such an access scheme, a channel is operated at about a Gb/s, and is allocated to a single user or is shared by several users via Time Division Multiplexing (TDM) or Sub-Carrier Multiplexing (SCM). The basic devices that are necessary for this mode of access are tunable optical filters, tunable transmitters, and optical amplifiers. While several hundred channels are theoretically possible with WDM, current optical technology limits the number of available wavelength channels to a few tens, typically less than 30.

The first WDM networks that were studied are classified as broadcast-and-select networks, in which the transmission from each network node is broadcast to all nodes. The desired signal is then extracted from a combination of signals at different wavelengths at the receiver. A great amount of research has been performed in studying broadcast-and-select networks using the wavelength-insensitive passive optical star coupler [5].

Despite their simplicity, broadcast-and-select networks have two inherent drawbacks. Firstly, the absence of wavelength selectivity within the network and the broadcast nature of the network prevent the concurrent use of any wavelength by more than one session. Secondly, since each signal is broadcast to all the network nodes, there is a power *splitting loss* [2]. The simplicity of broadcast-and-select networks makes them a very attractive alternative to conventional electronic networks in the local and metropolitan area settings. However, because of these deficiencies, they are not suitable for the wide area, much like the way electronic broadcast networks such as the Ethernet are limited to local-area-network (LAN) and metropolitan-area-network (MAN) environments.

The development of optical devices that perform wavelength-sensitive routing has given rise to the wavelength routing concept [6]. At present, switching high-speed connections via wavelength routing is a promising concept for a wide-area-network (WAN). Circuit-switched networks employing wavelength routing will be the focus of

study in this dissertation. In Chapter 2, we provide a brief introduction to wavelength routing and the importance of wavelength conversion in wavelength-routing (also called wavelength-routed) WDM networks, and discuss the organization of the rest of the dissertation.

Chapter 2

WAVELENGTH ROUTING

2.1 *Wavelength Routing*

As mentioned in Chapter 1, wavelength routing uses wavelength sensitive devices to route the optical signals in the network. Examples of such devices are grating multiplexers and demultiplexers. Unlike in broadcast networks in which a wavelength cannot be used by two or more concurrent sessions, wavelength reuse is achievable in wavelength-routing networks. Wavelengths can be reused as long as no two sessions that share the same link use the same wavelength at the same time. Wavelength reuse is a critical characteristic that makes networks employing wavelength routing very attractive in a WAN.

In a wavelength-routing network, the path of a signal is determined by the location of the signal transmitter, the wavelength on which it is transmitted, and the state of the network devices. An example of such a network is shown in Figure 2.1. There are three sessions that are in progress, one from node 1 to node 4 using wavelength λ_1 , another from node 2 to node 3 using wavelength λ_1 , and a third from node 1 to node 3 on λ_2 , without any conflict in wavelength usage. Note that λ_1 is used in two sessions simultaneously while the third session could not use λ_1 because of wavelength conflict. Also note that a single wavelength is used along the entire path of a session.

A wavelength-routing network may be classified as *configurable* or *fixed* depending on the capability of the network devices. In a passive network, the path of a signal is completely determined by its origin and the wavelength on which it is launched. Routing occurs by tuning the transmitter/receiver to the appropriate channel and

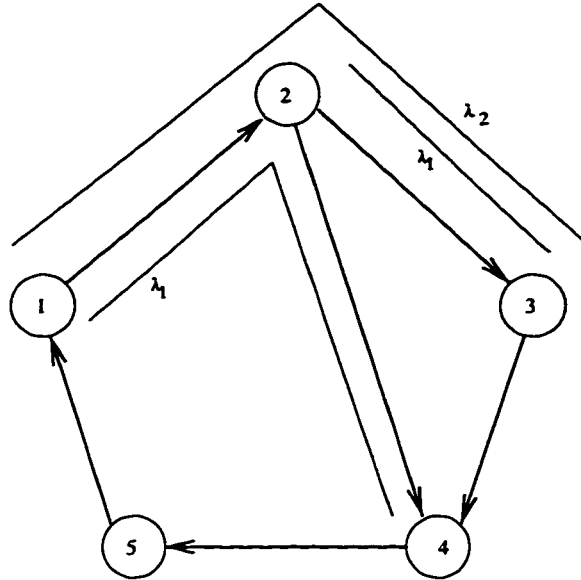


Figure 2.1: A wavelength-routing network.

then transmitting the information. On the other hand, reconfigurable network devices can be used for routing in a configurable network. Tunable transceivers are not required for all-to-all connectivity in such networks.

2.2 Routing Node Architectures

In this section, we review some of the possible architectures for a wavelength-routing node.

Passive Routing Node

Figure 2.2 [2] is an example of a passive routing node with three wavelengths per fiber. The signals on different wavelengths arriving at an input port are first demultiplexed by a grating demultiplexer, routed, and then multiplexed back on to a single fiber at the output. This architecture in which each input is connected to each output on exactly one wavelength is called a latin router [7]. Routing is performed

by transmitting at different wavelengths.

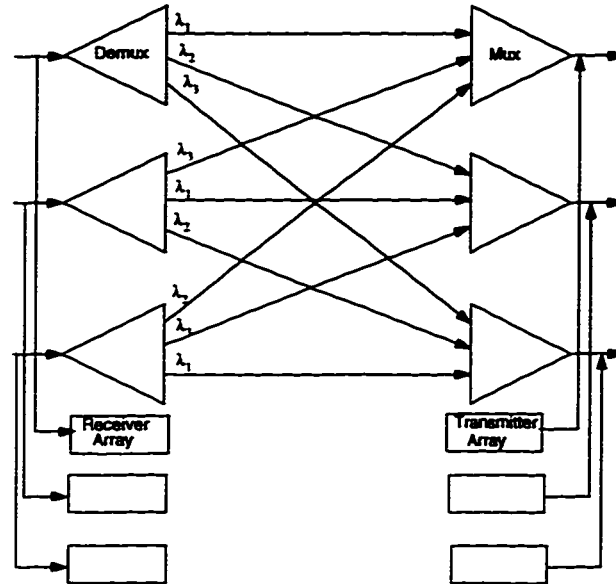


Figure 2.2: A passive wavelength-routing node.

Configurable Routing Node

The tuning requirements on transceivers and the number of wavelengths necessary for a given connectivity can be reduced by making the routing node dynamically reconfigurable [8, 9]. A possible architecture for such a routing node is given in Figure 2.3 [10].

In this network, the wavelengths on each fiber are first demultiplexed and each wavelength is connected to a photonic switch [11] dedicated to that wavelength. The various wavelengths coming out of the switches are then multiplexed back on a single fiber. In this architecture, $F N \times N$ optical space switches are used, where N is the number of ports and F is the number of wavelengths per fiber. This architecture allows two sessions emanating from the node to be on the same wavelength as long as they are on different output links. An architecture that greatly reduces the hardware

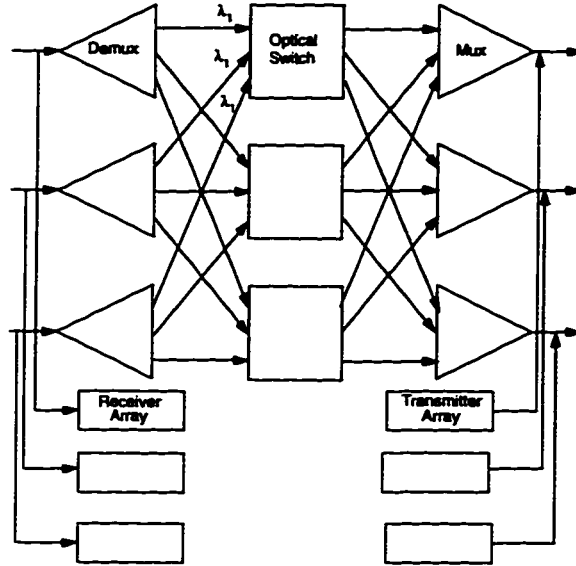


Figure 2.3: A configurable wavelength-routing node.

is possible if this facility is not required [10]. In this thesis, we assume that the network is configurable and also that each routing node has arrays of transmitters and receivers (one transmitter/receiver per wavelength per I/O port) as shown in Figure 2.3.

2.3 Wavelength Conversion

The discussion up to this point has assumed that a session uses the same wavelength along its entire path. Wavelength reuse can be improved by relaxing this wavelength continuity constraint as can be seen from the following example. Consider the network of Figure 2.1 and suppose the session (1, 4) uses λ_1 but the session (2, 3) is established on λ_2 . Then, with the wavelength continuity constraint, session (1, 3) will require a third wavelength. Without the continuity constraint, session (1, 3) can use λ_2 and λ_1 on its two links as shown in Figure 2.4, thereby saving λ_3 for a future session. This requires that node 2 have the capability to not only route the signal on λ_2 on link

(1, 2) to link (2, 3), but also to convert its wavelength to λ_1 . This capability is called wavelength conversion or wavelength translation.

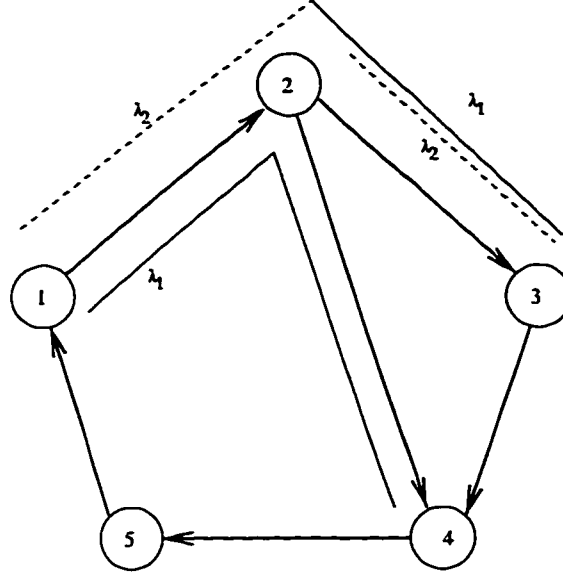


Figure 2.4: A wavelength-routing network with wavelength conversion.

A simple way to perform wavelength conversion is to convert the optical signal to electronic form and retransmit the information on another wavelength. This, however, reintroduces the electronic bottleneck and is not desirable. All-optical wavelength conversion techniques include (a) four-wave mixing in semiconductor optical amplifiers (SOAs) [12], (b) cross-gain modulation in SOAs [13] and (c) cross-phase modulation in SOAs [14].

Various architectures for wavelength-routing nodes with conversion capability are presented in [15]. In a node with full wavelength conversion, shown in Figure 2.5, any wavelength can be converted to any other wavelength since each wavelength on each link has its own dedicated converter. Each converter is capable of converting an arbitrary input wavelength to a fixed output wavelength. Note that the size of the switch now is $FN \times FN$ where F is the number of wavelengths and N is the number

of I/O ports.

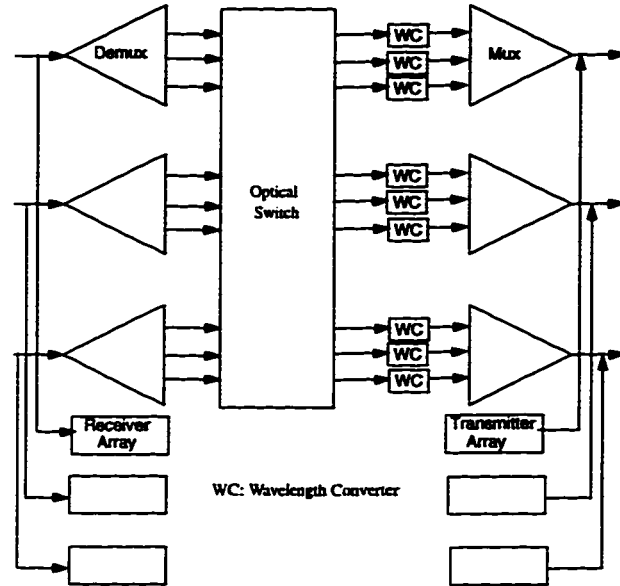


Figure 2.5: A routing node with full wavelength conversion.

A possibly more cost effective solution is to provide partial wavelength conversion at each node. Two possible structures for partial wavelength conversion are the following [15]: the share-per-node architecture of Figure 2.6 and the share-per-link architecture of Figure 2.7.

In the share-per-node architecture, the limited number of wavelength converters are collected in a single converter bank that can be accessed by all circuits passing through that node. An extra photonic switch is required to switch the circuits using the converter bank to their respective output links.

In the share-per-link architecture, there is a converter bank for each output link and the converters in a single bank are shared only between all the circuits using the corresponding output link. The amount of switching hardware is reduced in this architecture at the expense of reduced converter utilization.

Having presented an introduction to wavelength routing and pointed out the role

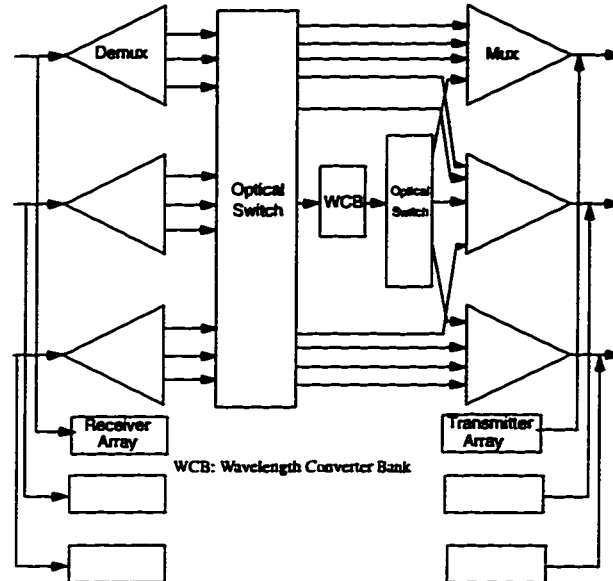


Figure 2.6: The share-per-node wavelength-convertible switch architecture.

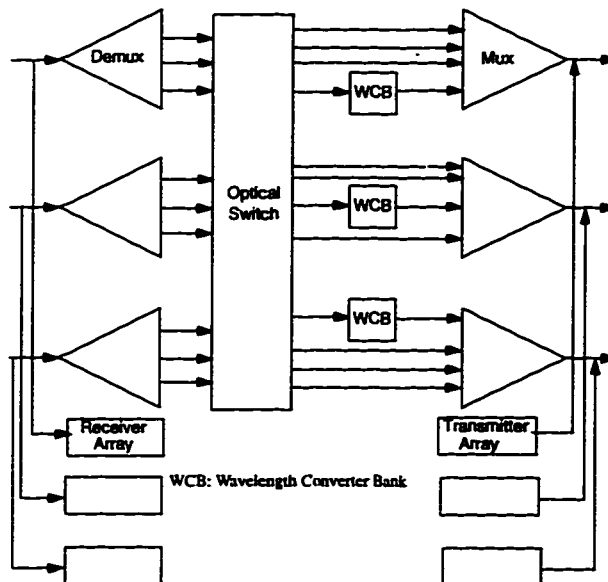


Figure 2.7: The share-per-link wavelength-convertible switch architecture.

of wavelength conversion in wavelength-routing networks, we now present the contributions and the outline of this dissertation.

2.4 Contributions and Outline of the Dissertation

The rest of the dissertation is organized as follows. All-optical wavelength conversion technology is not a commercially feasible technology at present. An important question that arises then is: How beneficial are wavelength converters in wavelength-routing networks? We address this question in Chapter 3. The benefits of wavelength conversion have been studied by assuming three types of traffic models: (a) static traffic, where all the connection requests are known ahead of time and the maximum supportable loads with zero, limited, and full wavelength conversion are compared [16], (b) dynamic traffic, where requests arrive and depart in a fashion such that the number of sessions on the most congested link does not exceed a certain constant [17], and (c) statistical dynamic traffic, where connection requests arrive and depart from the network in a random manner. We assume the third traffic model in this dissertation as we believe that this is more appropriate for a future network whose traffic cannot be predicted. Several performance models that estimate the benefits of wavelength conversion under statistical dynamic traffic exist [18, 19, 20, 21, 22]. All of these models have assumed that either there are no wavelength converters in the network or every node in the network is capable of full wavelength conversion. Because wavelength converters are expensive devices, a more practical scenario is one in which there is a limited number of wavelength converters in the network. We present a general model for analyzing the blocking probability of networks with *sparse wavelength conversion*, in which only a fraction of the network nodes is equipped with full wavelength conversion capability. The model is quite accurate for Poisson traffic even with sparsely connected topologies such as the ring for which other models do not perform well. The model's complexity is modest accounting for the fact that it

considers link load correlation which cannot be ignored in sparse topologies. As we will see in Chapter 3, the model can also be used to evaluate the performance of networks with all nodes having the share-per-node architecture, discussed earlier in this chapter. A related model to ours is the one presented in [23]. In that work, the authors assume that a wavelength can be converted to one of only a subset of wavelengths. However, that model is applicable only to dense topologies and does not take the dynamics of the traffic fully into account. The effect of non-Poissonian traffic is studied in Chapter 4 using a simpler version of the performance model in Chapter 3 and a simple non-Poisson traffic model called the Bernoulli-Poisson-Pascal (BPP) model.

We show in Chapter 3 that full wavelength conversion at all nodes may not be necessary for obtaining a desired level of performance, depending on the topology, number of wavelengths available, traffic, etc. A natural question then is the optimal placement of a given number of converters in a network. We present in Chapter 5 a dynamic programming solution to obtain the placement of a given number of converters along a path that minimizes the blocking probability. We also present exact placement results for bus and ring topologies.

In Chapter 6, we consider networks with no wavelength conversion. Routing and wavelength assignment (RWA), i.e., the selection of an appropriate path and wavelength, for an arriving connection request is an important issue in circuit-switched wavelength-routed networks. The wavelength assignment problem does not exist in networks with full wavelength conversion at every node. Even with static traffic with given routes, the problem of wavelength assignment to minimize the number of wavelengths used is NP-complete [24]. For statistical dynamic traffic, it is impossible to specify an exact RWA policy that minimizes the blocking probability because of the unpredictability of future traffic. Given the difficult nature of the problem, one has to resort to heuristics to obtain good solutions. In Chapter 6, we consider networks with fixed routing and focus on the wavelength assignment problem. We

present a heuristic that outperforms other known wavelength assignment algorithms, for Poisson traffic.

Conclusions and directions for future work are presented in Chapter 7.

Chapter 3

AN ANALYTICAL MODEL FOR SPARSE WAVELENGTH CONVERSION: POISSON TRAFFIC

The focus of this chapter is the blocking performance of wavelength-routing networks under Poisson input traffic. There has been considerable interest in obtaining the call blocking performance of wavelength-routing networks [18, 19, 21, 22, 23, 25]. The performance improvement with wavelength converters is of fundamental importance to quantify. This improvement depends on the topology of the network, the traffic demand, and the number of available wavelengths, among other factors. As the network becomes denser, one would expect the usefulness of converters to decrease, since the paths get shorter. In the limiting case with a link between every node pair, wavelength converters have no effect on the blocking performance, since all sessions are one-hop sessions¹. On the other hand, a sparsely connected network tends not to mix calls well and thus causes a load correlation in successive links. This reduces the usefulness of wavelength converters [18]. The benefits of conversion are thus largely dependent on which of the above two effects dominates.

The analytical models proposed in the literature have considered networks in which there are no converters and those in which all nodes have converters. All-optical wavelength converters are being prototyped in research laboratories [26], and are likely to remain costly devices. Therefore, a more practical situation is one in which wavelength conversion capability is available in a relatively small fraction of

¹ This assumes that the direct link is always used in this case. If alternate path routing is allowed, wavelength converters may still be of some benefit.

nodes. Our model generalizes the previous models in that we consider networks which have full wavelength conversion capability at an arbitrary fraction of the network nodes. We call a network that does not have full conversion capability at every node as one with *sparse wavelength conversion*.

The literature on the blocking analysis of networks with full or no wavelength conversion points out the difficulty in accurate accounting of load correlation. In [22], a model to compute the approximate blocking probability with Poisson traffic input is presented. As pointed out in [22], that model is inappropriate for networks with sparse topologies because it does not consider the correlation of wavelength use between successive links of a path. In [18], a model that takes into account this dependence is presented. However, it assumes that a wavelength is used on a link with a fixed probability independently of the other wavelengths, and, therefore, is not numerically accurate for Poisson traffic. A more accurate version of that model for Poisson traffic was presented in [25]. Another model presented in [19] considers Poisson input traffic and uses a Markov chain model with state-dependent arrival rates. It is very accurate, but the computation of blocking probabilities is very intensive, and the analysis is tractable only for networks with a few nodes.

In [27], we presented a model for sparse wavelength conversion that improves on the model presented in [22] by considering link load correlation to a certain extent. In this chapter, we describe a more accurate model for networks with sparse wavelength conversion [28].

We consider circuit-switched networks. Call requests arrive at a node according to a random point process and an optical circuit is established between the source and destination for the (random) duration of the call. When all nodes have wavelength converters, the situation is analogous to trunk switching in digital telephony [29] as a call arriving on one trunk (wavelength) can be switched to any outgoing trunk, so long as one is available. On the other hand, in a network without any wavelength converting nodes, a call arriving at a certain wavelength on an input fiber has to be

switched to an output fiber at the same wavelength. This requirement of wavelength continuity increases the probability of call blocking; to honor a call request, it is necessary that the *same* wavelength be free on all the links of the circuit.

In Section 3.1, we present a model for computing the blocking performance of networks without wavelength converters. The model takes into account the correlation of wavelength use on successive links. Section 3.2 extends this model to incorporate the effect of sparse wavelength conversion for an arbitrary topological connectivity. We then study the blocking performance of ring and mesh-torus networks as examples of sparse topologies in Sections 3.3 and 3.4 respectively, and the hypercube as an example of a dense topology in Section 3.5. We notice that there is a good match between the analytical and simulation results even when the topologies are sparse. The results of the study are counterintuitive and the effect of wavelength conversion on the blocking performance of a given network topology cannot be predicted *a priori* without an accurate analysis.

Wavelength-routing is most likely to be used in wide-area-networks (WANs) where a broadcast-and-select approach to switching ceases to be feasible [30]. The topology of such a network typically evolves into an irregular physical topology with arbitrary connection patterns. To model this practical situation, we consider random topologies in Section 3.6. We evaluate the ensemble average of the blocking probability and study the effects of connectivity and wavelength conversion. Our conclusions are presented in Section 3.7.

3.1 An Analytical Model for Path Blocking Performance

In this section, we develop an analytical model that has modest computational requirements and improves upon the previously proposed models of [18] and [22] by considering real-time input traffic and by incorporating the correlation between the wavelengths used on successive links of a multilink path. We first assume that there

is no wavelength conversion. In Section 3.2, we include the effect of wavelength conversion.

Model Assumptions

The following assumptions are used in our model:

1. Call requests arrive at each node according to a Poisson process with rate λ . Each call is equally likely to be destined to any of the remaining nodes.
2. Call holding time is exponentially distributed with mean $1/\mu$; the offered load per station² is $\rho = \lambda/\mu$.
3. The path used by a call is chosen according to a prespecified criterion (e.g., random selection of a shortest path), and does not depend on the state of the links that make up a path. The call is blocked if the chosen path cannot accommodate it. Alternate path routing is not allowed. This assumption is used to simplify the analysis. In practice, alternate routing will probably be used. Analyses for networks employing alternate routing are presented in [19, 31, 32].
4. The number of wavelengths, F , is the same on all links. Each node is capable of transmitting and receiving on any of the F wavelengths. Each call requires a full wavelength on each link it traverses.
5. Wavelengths are assigned to a session randomly from the set of free wavelengths on the associated path³.

² The words “node” and “station” are used interchangeably throughout this dissertation.

³ Other heuristic wavelength allocation strategies may provide better performance [32], but are considerably more difficult to analyze.

First, let us define a wavelength as “free” on a path if that wavelength is not used on any of the links constituting the path. A wavelength is “busy” on a path otherwise. The model in [22], called the *independence model* henceforth, assumes that the link loads are independent *and* that the wavelengths used on a link are uniformly distributed over the entire set of wavelengths, independently of all other links. These two assumptions result in an overestimation of the blocking probability, as shown in Section 3.3. The performance estimate is very crude (an inaccuracy of about two orders of magnitude) for sparsely connected networks such as the ring, and gets better with increasing connectivity. In this chapter, we propose a model that can be used to estimate the performance reasonably accurately even for ring topologies.

Notation

We define the following steady-state probabilities that will be used in obtaining the blocking probabilities.

- $Q(w_f) = \Pr\{w_f \text{ wavelengths are free on a link}\}.$
- $S(y_f|x_{pf}) = \Pr\{y_f \text{ wavelengths are free on a link of a path} \mid x_{pf} \text{ wavelengths are free on the previous link of the path}\}.$
- $U(z_c|y_f, x_{pf}) = \Pr\{z_c \text{ calls (wavelengths) continue to the current link from the previous link of a path} \mid x_{pf} \text{ wavelengths are free on the previous link, and } y_f \text{ wavelengths are free on the current link}\}.$
- $R(n_f|x_{ff}, y_f, z_c) = \Pr\{n_f \text{ wavelengths are free on a two-hop path} \mid x_{ff} \text{ wavelengths are free on the first hop of the path, } y_f \text{ wavelengths are free on the second hop, and } z_c \text{ calls continue from the first to the second hop}\}.$
- $T^{(l)}(n_f, y_f) = \Pr\{n_f \text{ wavelengths are free on an } l\text{-hop path and } y_f \text{ wavelengths are free on hop } l\}.$

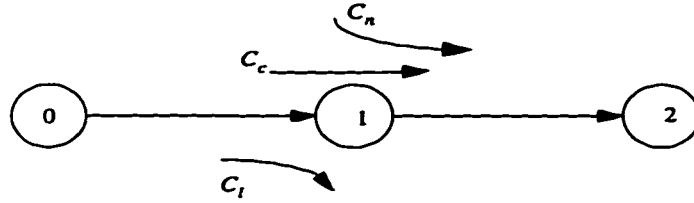


Figure 3.1: Calls arriving and leaving on a two-hop path.

- $p_l = \Pr\{\text{an } l\text{-hop path is chosen for routing}\}$.

The Conditional Free Wavelength Distribution

In addition to the assumptions stated in Section 3.1, we assume the following in the analytical model. The load on link i of a path given the loads on links $1, 2, \dots, i-1$, depends only on the load on link $i-1$. We therefore call our analytical model as the (Markovian) *correlation model*. This analysis differs from the one presented in [27] which neglected link load correlation. We start by considering a two-hop path and deriving the conditional free wavelength distribution on the path. In Section 3.1, we extend the analysis to determine the blocking probability on a path of arbitrary hop length. To model the correlation between the loads on the two links, we employ a three-dimensional Markov chain as follows. Referring to Figure 3.1, let C_l be the number of calls that enter the path at node 0 and leave at node 1, let C_c be the number of calls that enter the path at node 0 and continue on to the second link, and let C_n be the number of calls that enter the path at node 1. Therefore, the number of calls that use the first link is $C_l + C_c$ and the number of calls that use the second link is $C_c + C_n$.

Since the number of calls on a link cannot exceed the total number of available wavelengths, F , we have $C_l + C_c \leq F$, and $C_c + C_n \leq F$. Suppose the arrival rate of calls that enter at node 0 and leave at node 1 is λ_e , and the arrival rate of calls that enter at node 0 and continue on to node 2 is λ_c . Let the corresponding Erlang

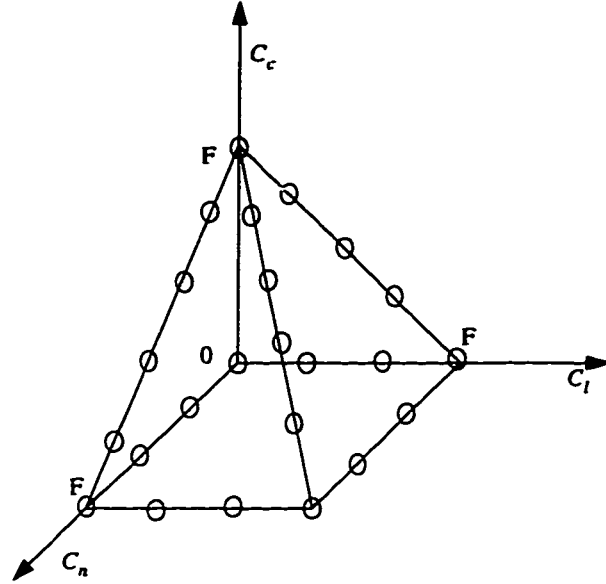


Figure 3.2: The state-space of the three-dimensional Markov chain.

loads be denoted by $\rho_e = \lambda_e/\mu$ and $\rho_c = \lambda_c/\mu$, where $1/\mu$ is the expected value of the exponentially distributed call holding time. By the assumption of uniform traffic distribution, the arrival rate of calls that enter the path at node 1 is the same as the arrival rate of calls that leave the path at node 1. C_l , C_c , and C_n can therefore be characterized by a three-dimensional Markov chain, with the state space as shown in Figure 3.2. Each state is represented by an integer triplet (c_l, c_c, c_n) (the circles in Figure 3.2). λ_e and λ_c actually change with the state of this Markov chain because of blocking on other links of the network. This situation is modeled in [19] for the link-load independence case. Modeling it in our case would considerably complicate the model (in particular, the steady-state probabilities of (3.1) are very hard to obtain numerically), and we show later that the loss in accuracy due to our assumption of constant λ_e and λ_c is not significant.

We can now determine the steady-state probability of state (c_l, c_c, c_n) as [33],

$$\pi(c_l, c_c, c_n) = \frac{\frac{\rho_c^{c_l} \rho_c^{c_c} \rho_c^{c_n}}{c_l! c_c! c_n!}}{\sum_{j=0}^F \sum_{i=0}^{F-j} \sum_{k=0}^{F-j} \frac{\rho_c^i \rho_c^j \rho_c^k}{i! j! k!}}, \quad 0 \leq c_l + c_c \leq F, \quad 0 \leq c_c + c_n \leq F. \quad (3.1)$$

In the absence of wavelength converters, a call is blocked on a two-hop path if there is no free wavelength on the path. Now we compute the free wavelength distribution on a two-hop path. The conditional probability that n_f wavelengths are free on a two-hop path given x_{ff} wavelengths are free on the first link, y_f wavelengths are free on the second link, and z_c calls continue from the first link to the second link, is obtained by simple enumeration as

$$R(n_f | x_{ff}, y_f, z_c) = \frac{\binom{x_{ff}}{n_f} \binom{F-x_{ff}-z_c}{y_f-n_f}}{\binom{F-z_c}{y_f}}$$

for $\min(x_{ff}, y_f) \geq n_f \geq \max(0, x_{ff} + y_f + z_c - F)$, and is 0 otherwise. We have assumed here that n_f depends only on the numbers x_{ff} , y_f , and z_c , and not on how the system reached the state (x_{ff}, y_f, z_c) . Effectively, we have assumed that the calls are arranged on the two links such that all wavelength configurations that are possible when the system is in state (x_{ff}, y_f, z_c) are equiprobable. Modeling the exact situation here is very difficult.

Recalling our notation in Section 3.1, the probabilities $U(z_c | y_f, x_{pf})$, $S(y_f | x_{pf})$, and $Q(w_f)$ can be derived using Figure 3.2 as follows:

$$\begin{aligned} U(z_c | y_f, x_{pf}) &= P(C_c = z_c \mid C_n + C_c = F - y_f, C_l + C_c = F - x_{pf}) \\ &= \frac{\pi(F - x_{pf} - z_c, z_c, F - y_f - z_c)}{\sum_{x_c=0}^{\min(F-x_{pf}, F-y_f)} \pi(F - x_{pf} - x_c, x_c, F - y_f - x_c)}, \end{aligned}$$

$$S(y_f | x_{pf}) = P(C_n + C_c = F - y_f \mid C_l + C_c = F - x_{pf})$$

$$= \frac{\sum_{x_c=0}^{\min(F-x_{pf}, F-y_f)} \pi(F-x_{pf}-x_c, x_c, F-y_f-x_c)}{\sum_{x_c=0}^{F-x_{pf}} \sum_{x_n=0}^{F-x_c} \pi(F-x_{pf}-x_c, x_c, x_n)},$$

and

$$\begin{aligned} Q(w_f) &= P(C_l + C_c = F - w_f) \\ &= \sum_{x_c=0}^{F-w_f} \sum_{x_n=0}^{F-x_c} \pi(F-w_f-x_c, x_c, x_n). \end{aligned} \quad (3.2)$$

Next, we use the conditional free wavelength distribution to obtain an expression for the blocking probability (in the absence of wavelength conversion) on a path of arbitrary hop-length.

Blocking on a Multihop Path

Consider an l -hop path. A call is blocked if there is no free wavelength on the path at the time of call arrival. Suppose we know the joint probability, $T^{(l-1)}(x_{ff}, x_{pf})$, that there are x_{ff} free wavelengths on an $(l-1)$ -hop path and $x_{pf} (\geq x_{ff})$ wavelengths are free on the last hop of that path. Because of our assumption that the load on the l th hop is dependent only on the load on the $(l-1)$ th hop, the probability of blocking on the l -hop path can be computed using the results for a two-hop path derived in Section 3.1. This is done by viewing the first $l-1$ hops as the first hop and the l th hop as the second hop of a two-hop path. To complete the recursion, we need to determine the joint probability of n_f free wavelengths on the l -hop path and y_f free wavelengths on the l th hop.

Using the chain rule of probability, we have

$$\begin{aligned} T^{(l)}(n_f, y_f) &= \sum_{x_{pf}=0}^F \sum_{z_c=0}^{\min(F-x_{pf}, F-y_f)} \sum_{x_{ff}=0}^{x_{pf}} \\ &\quad R(n_f | x_{ff}, z_c, y_f) U(z_c | y_f, x_{pf}) S(y_f | x_{pf}) T^{(l-1)}(x_{ff}, x_{pf}). \end{aligned} \quad (3.3)$$

In writing the above, we have used the following facts:

- (i) the number of free wavelengths on the l -hop path is not dependent on the number of free wavelengths on hop $l - 1$ when the number of free wavelengths on the path consisting of the first $l - 1$ hops is given,
- (ii) the number of calls that continue from the first to the second link of a two-hop path depends only on the number of free wavelengths on each of the two hops, and
- (iii) the number of free wavelengths on the second link of a two-hop path is dependent only on the number of free wavelengths on the first link.

The starting point of the recursion, $T^{(1)}(n_f, y_f)$, is zero when $n_f \neq y_f$, and is equal to the probability of having n_f free wavelengths on a link, $Q(n_f)$, given by (3.2), when $n_f = y_f$.

The probability of blocking on an l -hop path is simply the probability of finding no wavelength free on the l -hop path, and is therefore given by $\sum_{y_f=0}^F T^{(l)}(0, y_f)$. The network-wide blocking probability in the absence of converters is then computed as

$$P_b = \sum_{l=1}^{N-1} \sum_{y_f=0}^F T^{(l)}(0, y_f) p_l$$

for a network of N nodes.

Estimation of Parameters

The analysis in the last two sections can be used to compute the call blocking probability in a network without wavelength converters. This analysis assumes that the hop-length distribution, p_l , and the arrival rates of calls at a link that continue on to the next link of a path and of those that do not, λ_c and λ_e respectively, are known. The hop-length distribution is a function of the topology and the routing algorithm and is easily determined for most regular topologies with shortest-path routing, as seen in the following sections.

Typically, the traffic in the network is specified in terms of the set of offered loads between station pairs. The call arrival rates at links have to be estimated from the arrival rates of calls to nodes. The complication in estimating the link arrival rates is that the entire offered load is not carried by the network because of call blocking. The probability of blocking is, in turn, dependent on the arrival rate to the links. This leads to a system of coupled non-linear equations called the *Erlang map* [29]. While solving the Erlang map leads to a more accurate computation of blocking probabilities, the effect of blocking probabilities on the carried load can be neglected, especially when these probabilities are small. We take this approach to keep the analysis simple.

When the blocking probabilities are small, the link arrival rates, λ_e and λ_c , can be computed as follows. Consider a network with N nodes. Let the total arrival rate of calls at a link be γ ($= \lambda_e + \lambda_c$). The sum of arrival rates of calls at all links in the network is γL , where L is the number of links in the network. Since a call uses a path of expected length $\bar{H} = \sum_{l=1}^{N-1} l p_l$, the sum of arrival rates at all links in the network is also $N\lambda\bar{H}$ where λ is the call arrival rate at a node. Thus,

$$\gamma = \frac{N\lambda\bar{H}}{L}. \quad (3.4)$$

(For networks that are asymmetric, the arrival rates at all links may not be the same even if the node arrival rates are the same. In such cases, γ is the arrival rate averaged over all the links in the network.)

Having obtained the total arrival rate per link, we now estimate the arrival rates λ_e and λ_c . A plausible estimate for the probability of a call leaving the network at a given node⁴ is $1/\bar{H}$. Now, consider an intermediate node of a path and define an *exit link* of this node as an outgoing link that does not return to the previous node of the path. Suppose there are k exit links per node. We assume that if a call does not leave

⁴The implicit assumption here is that the hop-length of a call is geometrically distributed with mean $\bar{H}$.

the network at the node, it chooses one of the k exit links arbitrarily⁵. Therefore, given a path and a link of the path, the arrival rate of calls that continue on to the next link of the path can be estimated as

$$\lambda_c = \gamma \frac{1 - 1/\overline{H}}{k}. \quad (3.5)$$

Calls that do not continue on to the next link of the path either leave the network or continue on a different link from the given node. The arrival rate of such calls at the given link is simply given as

$$\lambda_e = \gamma - \lambda_c. \quad (3.6)$$

It is worthwhile observing at this point that the correlation model we have just presented subsumes the independence model that assumes that link loads are independent. When λ_c is set to 0 and λ_e is set to γ , our model greatly simplifies and reduces to the model presented in [22]. The model in [22] actually solves the Erlang map and iteratively computes the blocking probabilities. We, however, ignore the effect of reduced load due to blocking. Incorporating the effect in our model would not be hard but would complicate the presentation of the model unnecessarily. The independence model is shown to be inaccurate for sparse topologies in Section 3.3, and, more importantly, the correlation model is shown to be quite accurate.

3.2 Sparse Wavelength Conversion

In the previous section, we assumed that wavelength converters are absent in the network. Wavelength converters improve the blocking performance by allowing a circuit to be established as long as *some* wavelength is available on each link of the desired path. To enhance the blocking performance of the network, wavelength converters are placed at some of the nodes. All the previous analyses in the literature

⁵ For irregular topologies where the number of exit links may be different for different nodes, k could be taken as the average number of exit links per node.

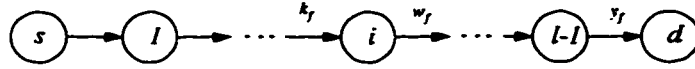


Figure 3.3: An l -hop path with nodes numbered $s, 1, 2, \dots, l-1, d$.

have considered networks without any converters or networks with converters at every node. An interesting design alternative that has not been considered previously is one in which wavelength conversion is available in a subset of network nodes to achieve a balance between cost and performance.

We model a network with sparse wavelength conversion by assuming that a node is capable of wavelength conversion with probability q independently of the other nodes. q is called the *conversion density* of the network. The number of converter nodes in an N -node network is thus binomially distributed with an average of Nq converters, and the blocking performance we obtain is the ensemble average of the blocking probability over this distribution. This probabilistic approach enables a single parameter characterization of the wavelength conversion density and eliminates the unpleasant task of evaluating the performance for each number and placement of converters.

Note that we assume that a node can either convert any set of wavelengths to any other, or cannot convert any wavelength. However, our analysis also applies to the case of limited wavelength conversion per node discussed in [15]. In particular, we predict the ensemble average blocking performance of the share-per-node architecture of [15] where the number of conversions at each node is binomially distributed with a mean of FDq , D being the number of links per node.

Consider an l -hop path in a network with conversion density q and let the nodes on the path be numbered $s, 1, 2, \dots, l-1, d$, as shown in Figure 3.3.

Let the call blocking probability on the path be denoted by $P_b^{(l)}(q)$. We recursively compute the call blocking probabilities on paths of different hop lengths. The idea

behind the recursion is as follows. Suppose node i is the last converter on the l -hop path. A call is not blocked on the path if (a) the call is not blocked on the first i hops of the path, *and* (b) there is a wavelength that is free on the last $l - i$ hops of the path. These two events, however, are not independent because the probability of blocking on the last $l - i$ hops depends on the number of free wavelengths on hop $(i + 1)$. This depends on the number of free wavelengths on hop i which, in turn, is dependent on whether the call is blocked on the first i hops or not. To analyze this situation exactly, we introduce the following probabilities. Let

- $V^{(l)}(n_f, y_f | w_f) = \Pr\{n_f \text{ wavelengths are free on an } l\text{-hop path and } y_f \text{ wavelengths are free on hop } l \mid w_f \text{ wavelengths free on the first hop of the path and no converters along the path}\}.$
- $W^{(l)}(y_f | w_f) = \Pr\{y_f \text{ wavelengths are free on hop } l \text{ of an } l\text{-hop path} \mid w_f \text{ wavelengths free on the first hop of the path and no converters along the path}\}.$
- $P_b^{(l)}(q, y_f) = \Pr\{\text{a call is blocked on an } l\text{-hop path and } y_f \text{ wavelengths free on hop } l, \text{ when the conversion density is } q\}.$

(3.3) can then be modified to give

$$V^{(l)}(n_f, y_f | w_f) = \sum_{x_{pf}=0}^F \sum_{z_c=0}^Z \sum_{x_{ff}=0}^{x_{pf}} R(n_f | x_{ff}, z_c, y_f) \times U(z_c | y_f, x_{pf}) S(y_f | x_{pf}) V^{(l-1)}(x_{ff}, x_{pf} | w_f), \quad (3.7)$$

for $l = 2, 3, \dots, N - 1$, where $Z = \min(F - x_{pf}, F - y_f)$. For a one-hop path,

$$V^{(1)}(n_f, y_f | w_f) = \begin{cases} 1, & \text{if } n_f = y_f = w_f, \\ 0, & \text{otherwise.} \end{cases}$$

Furthermore, by summing over all possible values of n_f in (3.7), we obtain

$$W^{(l)}(y_f | w_f) = \sum_{n_f=0}^{y_f} V^{(l)}(n_f, y_f | w_f).$$

The range of the variable n_f in the above equation is only up to y_f because the number of free wavelengths on an l -hop path without converters cannot be higher than the number of free wavelengths on any hop, in particular, hop l .

On a one-hop path, a call is blocked if and only if there is no free wavelength on the hop. Therefore,

$$P_b^{(1)}(q, y_f) = \begin{cases} Q(0), & \text{if } y_f = 0, \\ 0, & \text{otherwise,} \end{cases}$$

where $Q(w_f)$ is the probability that w_f wavelengths are free on a link, as given by (3.2).

For $l \geq 2$, the joint probability, $P_b^{(l)}(q, y_f)$, can be computed recursively by conditioning on the disjoint events that node i is the last converter on the given path of l hops, $i = 1, 2, \dots, l-1$. Thus,

$$\begin{aligned} P_b^{(l)}(q, y_f) &= P(\text{Blocking and } y_f \text{ free on last hop} \mid \text{no converters in path}) \\ &\quad \cdot P(\text{no converters in the path}) \\ &+ \sum_{i=1}^{l-1} P(\text{Blocking and } y_f \text{ free on last hop} \mid \text{node } i \text{ last converter}) \\ &\quad \cdot P(\text{node } i \text{ last converter}). \end{aligned} \quad (3.8)$$

When there are no converters, the joint probability of a call being blocked and y_f wavelengths being free on the last hop of an l -hop path is given by $T^{(l)}(0, y_f)$ (see (3.3)). The first term in (3.8) is thus $T^{(l)}(0, y_f)(1 - q)^{l-1}$. The joint probability, $Y^{(l)}(i, q, y_f)$, of a call being blocked and y_f wavelengths being free on the last hop of an l -hop path when node i is the last converter is obtained as follows.

$Y^{(l)}(i, q, y_f)$ can be written as

$$\begin{aligned} Y^{(l)}(i, q, y_f) &= \Pr\{y_f \text{ wavelengths free on hop } l \mid \text{node } i \text{ last converter}\} - \\ &\Pr\{\text{No blocking on first } i \text{ and on last } l-i \text{ hops and } y_f \text{ free wavelengths on hop } l \mid \\ &\text{node } i \text{ last converter}\}. \end{aligned}$$

The probability that y_f wavelengths are free on hop l is $Q(y_f)$ and is statistically

independent of conversion density. Referring to Figure 3.3,

$$\Pr\{\text{No blocking on first } i \text{ and on last } l-i \text{ hops and } y_f \text{ free wavelengths on hop } l \mid \text{node } i \text{ last converter}\} = \sum_{k_f=0}^F \sum_{w_f=0}^F P(A)P(B)P(C)$$

where the events A , B , and C are as defined below.

A : Call is not blocked on the first i hops and k_f wavelengths are free on hop i .

B : w_f wavelengths are free on hop $i+1$ given k_f wavelengths are free on hop i .

C : Call is not blocked on the last $l-i$ hops and y_f wavelengths are free on hop l given that w_f wavelengths are free on hop $i+1$ and there are no converters along the $(l-i)$ -hop path.

Now,

$$\begin{aligned} P(A) &= Q(k_f) - P_b^{(i)}(q, k_f), \\ P(B) &= S(w_f | k_f), \text{ and} \\ P(C) &= W^{(l-i)}(y_f | w_f) - V^{(l-i)}(0, y_f | w_f). \end{aligned}$$

Substituting for $P(A)$, $P(B)$, and $P(C)$, we can now write

$$\begin{aligned} Y^{(l)}(i, q, y_f) &= Q(y_f) - \\ &\sum_{k_f=0}^F \left(Q(k_f) - P_b^{(i)}(q, k_f) \right) \sum_{w_f=0}^F S(w_f | k_f) \left(W^{(l-i)}(y_f | w_f) - V^{(l-i)}(0, y_f | w_f) \right). \end{aligned} \quad (3.9)$$

The resulting $P_b^{(l)}(q, y_f)$ is given by

$$P_b^{(l)}(q, y_f) = T^{(l)}(0, y_f)(1-q)^{l-1} + \sum_{i=1}^{l-1} Y(i, q, y_f)q(1-q)^{(l-i-1)}. \quad (3.10)$$

The network call blocking probability when the conversion density is q is then given by

$$P_b(q) = \sum_{l=1}^{N-1} \sum_{y_f=0}^F P_b^{(l)}(q, y_f) p_l.$$

Given a network topology, we first determine the hop-length distribution, p_l , and the number of exit links, k . Then, (3.4), (3.5), and (3.6) are used to obtain the call arrival rates. The blocking performance can then be studied by using the analysis

presented in Section 3.1 and this section. In the next four sections, we evaluate the call blocking performance for the ring, the mesh-torus, the hypercube, and random topologies.

3.3 Blocking in a Ring Network

In this section, we apply the above analysis to a unidirectional ring network. The ring is the most sparsely connected network with a given number of nodes⁶ and we are interested in finding the effect of conversion density on such a network.

To verify the accuracy of the proposed analytical framework, we performed a simulation study of the call blocking performance in ring networks of different sizes using the stated assumptions in Section 3.1. Each data point in the simulations was obtained using 10^6 call arrivals. We simulated the no-converter ($q = 0$) and full-converter ($q = 1$) cases and obtained the blocking probability to compare with the analytical results.

For an N -node unidirectional ring, the probability that an l -hop path is used for routing a session is, $p_l = \frac{1}{N-1}$ for $1 \leq l \leq N-1$, and k is 1. First, we plot the call blocking probability against the load per station for a 100-node network when the number of wavelengths per fiber are 5 and 20 in Figures 3.4(a) and 3.4(b), respectively. Analytical and simulation results are plotted for the no converter case ($q = 0$), and the full-converter case ($q = 1$). In both cases, the converter case curves lie below the corresponding no-converter case curves. The reasonably close match between the analytical and simulation results and the fact that the analytical results follow the trend of the simulation results indicate that the model is adequate in analytically predicting the performance in ring topologies. For comparison, we also plot the analytical results when $\lambda_e = \lambda \bar{H}$ and $\lambda_c = 0$ (independence model). It is seen

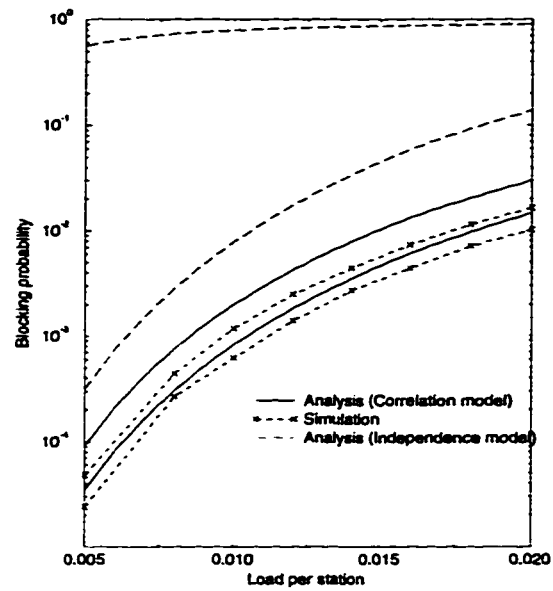
⁶ A star network is equally sparse in terms of the number of links, but has a lower average hop-length.

that the independence model severely overestimates the blocking probability⁷. We observe that wavelength converters are more useful when the number of wavelengths per fiber is larger and the load is lower. This can be explained as follows. When the number of wavelengths is larger, blocking occurs primarily not due to a lack of resources (wavelengths) but due to the inability of the network to use those resources efficiently in the absence of conversion. Thus, converters are more useful when the number of wavelengths is larger. Under heavy loads, blocking occurs primarily due to a lack of sufficient number of wavelengths and the presence of converters does not have as much effect as at lighter loads.

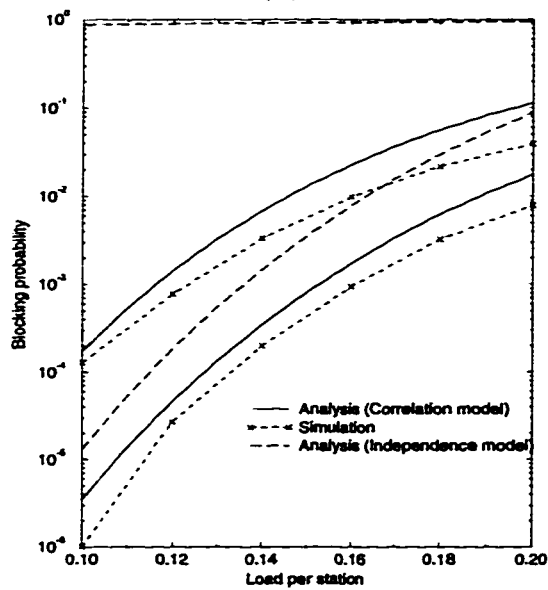
Figure 3.5 shows how the analytically obtained blocking probability changes with wavelength conversion density, q , for several values of F for a 20-node (solid lines) and a 100-node (dashed lines) ring network when the network load is 2 Erlangs (a load of 0.1 per station for the 20-node network and a load of 0.02 per station for the 100-node network). Two important observations can be made from these curves. Firstly, when the network has more nodes and/or the number of wavelengths is higher, the blocking probability drops fast initially with conversion density and then rapidly levels off at a certain point. Secondly, the density at which the performance begins to level off increases marginally as the number of wavelengths increases. In contrast, when the network is smaller, the blocking probability decreases more gradually. In either case, the decrease in blocking due to conversion density is almost insignificant (relative to the performance improvement we will see later in a mesh-torus). This leads to the observation that adding converters is not the solution to increasing the performance in very sparse networks.

We have also simulated the performance of a ring network with a fixed number Nq of converters which are randomly placed over the ring. We do not show the results of this experiment in order not to impair the clarity of Figure 3.5. These results indicate

⁷ The blocking probabilities are somewhat overestimated in all the curves because of the slight decrease in carried load due to blocking, which we have ignored.



(a)



(b)

Figure 3.4: The blocking probability vs. the load per station for a 100-node ring network. (a) $F = 5$, and (b) $F = 20$.

that the ensemble average is a very good approximation for this scenario.

Finally, we plot in Figure 3.6 the number of stations that can be supported for a blocking probability of 10^{-3} against the conversion density, for different loads. When the load per station is high (0.05 Erlangs), conversion density has very little effect on the number of stations that can be supported. This is again due to the limitation of the resources and not due to lack of efficient utilization of those resources. When the load becomes lighter (0.01 Erlangs), wavelength conversion helps initially in increasing the utilization but the performance begins to level off when the limit on resources is approached.

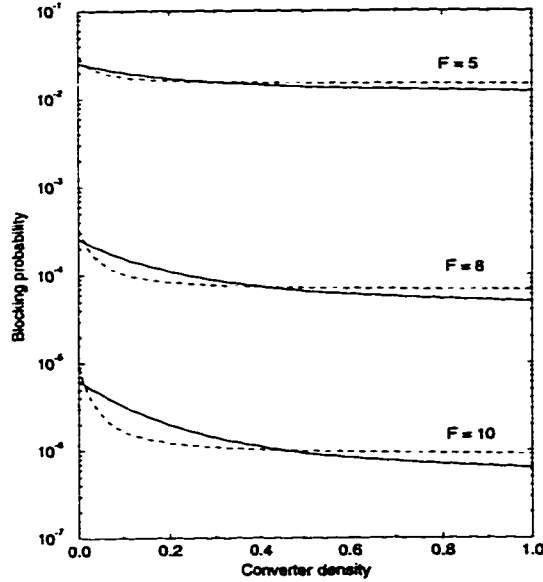


Figure 3.5: The blocking probability vs. the converter density for a ring network with a total network load of 2 Erlangs. Solid lines: $N = 20$ and $\rho = 0.1$, Dashed lines: $N = 100$ and $\rho = 0.02$.

3.4 Blocking in a Mesh-Torus Network

We next consider a bidirectional $M \times M$ mesh-torus network with $N = M^2$ nodes such as shown in Figure 3.7. A mesh-torus network is more connected than a ring

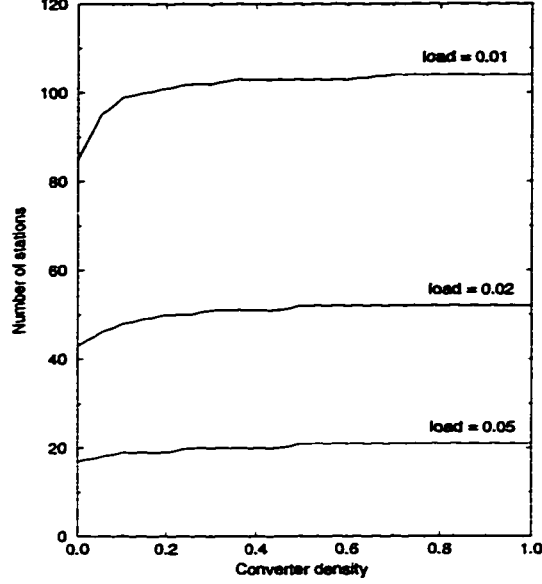


Figure 3.6: The converter density vs. the number of stations that can be supported for a ring network. $F = 5$, $P_b = 10^{-3}$.

but the average-hop length is still large. For simplicity in computing the hop-length distribution, we consider only odd values of M . We assume static routing where the route for a connection is randomly chosen from one of the many shortest path routes available. With this routing scheme, from any node there are $4l$ nodes at distance l if $1 \leq l \leq \frac{M-1}{2}$, and $4(M-l)$ nodes at distance l if $\frac{M-1}{2} < l \leq M-1$. Therefore, we have

$$p_l = \begin{cases} \frac{4l}{M^2-1}, & 1 \leq l \leq \frac{M-1}{2}, \\ \frac{4(M-l)}{M^2-1}, & \frac{M-1}{2} < l \leq M-1, \end{cases} \quad (3.11)$$

and the average hop-length is $\Theta(M)$. The number of exit links per node, k , is 3. The blocking performance is then analyzed by using the results of Sections 3.1 and 3.2.

In Figure 3.8, we plot the simulation and analytical (independence and correlation models) results for a 101×101 mesh-torus network with 5 wavelengths per fiber for the no-converter ($q = 0$) and full-converter ($q = 1$) cases. The results of our analytical model and the simulation results match very closely. Simulation results

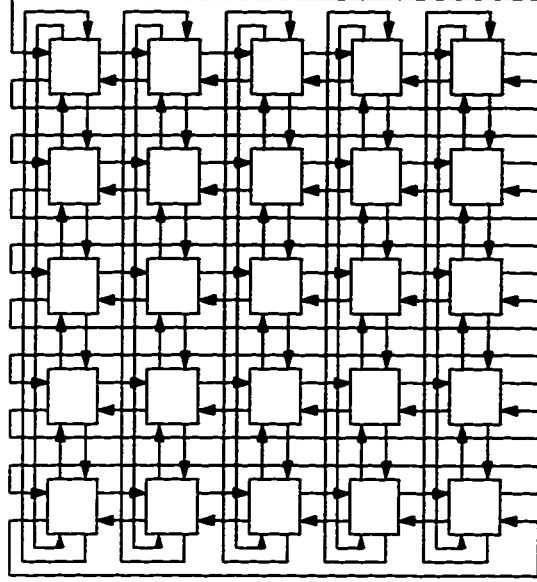


Figure 3.7: A 5×5 bidirectional mesh-torus network.

are reproduced from [22] with permission from the authors. Shortest path $X - Y$ routing is used in [22]. That routing algorithm yields the same distribution (3.11) for path lengths as random shortest path routing does. The independence model is less accurate than our model but not significantly so, indicating that the load correlation between successive links is very high in sparse networks and decreases as the network becomes more densely connected. We observe from the figure the tremendous improvement in performance with wavelength conversion, unlike in the case of the ring.

In Figures 3.9(a) and 3.9(b), we show the effect of conversion density on the analytically obtained blocking probability for a 11×11 and a 101×101 bidirectional mesh-torus network, respectively. As in the ring, we observe that conversion helps more when there are more wavelengths per fiber (the blocking probability drops more steeply with q as the number of wavelengths increases). Furthermore, the advantages of wavelength conversion are much higher in a larger network, and as the network size

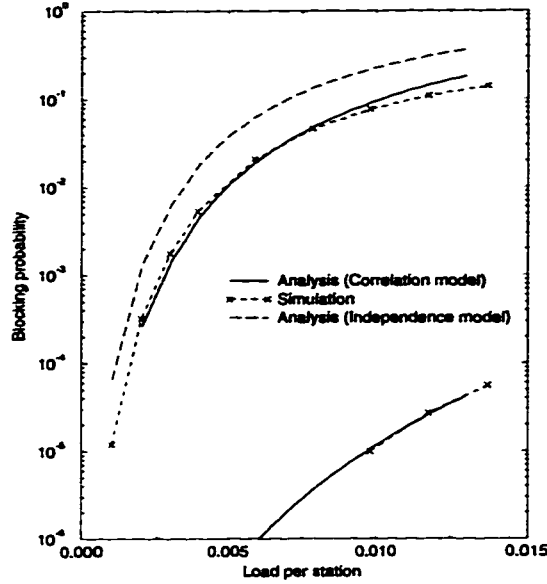


Figure 3.8: The blocking probability vs. the load per station for a 101×101 mesh-torus network with 5 wavelengths per fiber.

increases, performance increases dramatically initially with conversion density. For example, we observe from Figure 3.9(b) that, when $F = 5$, the blocking probability drops from 10^{-1} to 10^{-3} as the conversion density increases from 0 to about 0.2, and then decreases more gradually. When $F = 8$, a decrease in blocking probability of two orders of magnitude occurs as q increases from 0 to about 0.1. This not only suggests that having a converter at every node may be unnecessary to achieve a certain performance but also that the proportion of converter nodes required is a function of the size of the network and the number of wavelengths per fiber. However, unlike in the ring, when there are a large number of wavelengths, the performance does not level off sharply with increasing conversion density. This suggests that the utilization of converters is affected more by the large hop-lengths in the mesh than by the load correlation that dominates in the ring.

Figure 3.10 shows how the number of stations that can be supported increases with the conversion density. At lighter loads, the advantages of conversion are much

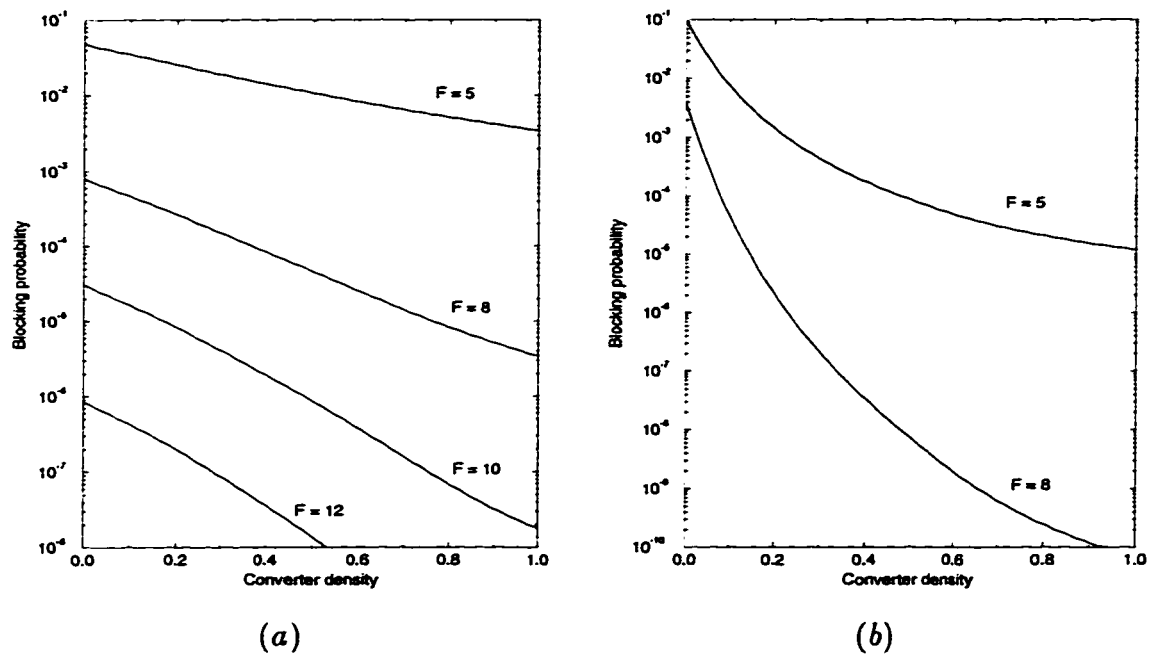


Figure 3.9: The blocking probability vs. conversion density for (a) a 11×11 mesh-torus with $\rho = 0.5$, and (b) a 101×101 mesh-torus with $\rho = 0.01$.

more significant than in the case of the ring. (Note that the square root of the number of stations is plotted on the Y -axis.) This is due to the better mixing of traffic in a mesh-torus although the paths are shorter than those in the ring.

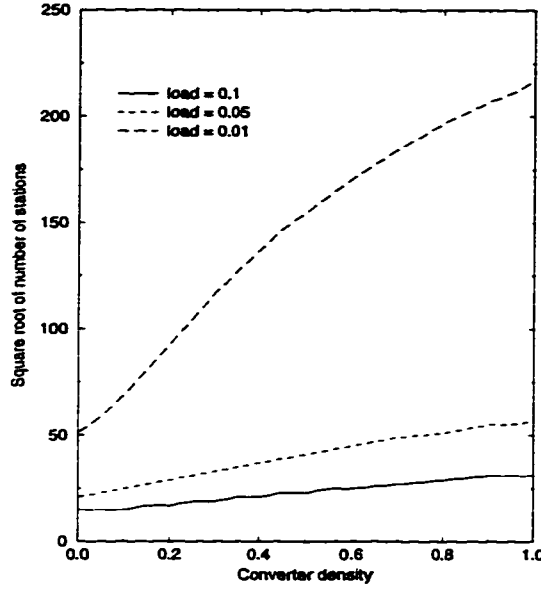


Figure 3.10: The converter density vs. the square root of the number of stations that can be supported for a bidirectional mesh-torus network. $F = 5$, $P_b = 10^{-3}$.

3.5 Blocking in a Hypercube Network

We next analyze a well-connected network, the binary hypercube (called hypercube henceforth). An $N(= 2^n)$ -node hypercube network has n outgoing links per node. We assume that one of the shortest paths between the source and destination nodes is chosen for routing a session. Using this routing scheme, we have

$$p_l = \frac{1}{N-1} \binom{n}{l}, \quad 1 \leq l \leq n,$$

and the average hop-length is $\Theta(n)$. The number of exit links per node, $k = n - 1$. The blocking performance is then analyzed using the results of Sections 3.1 and 3.2.

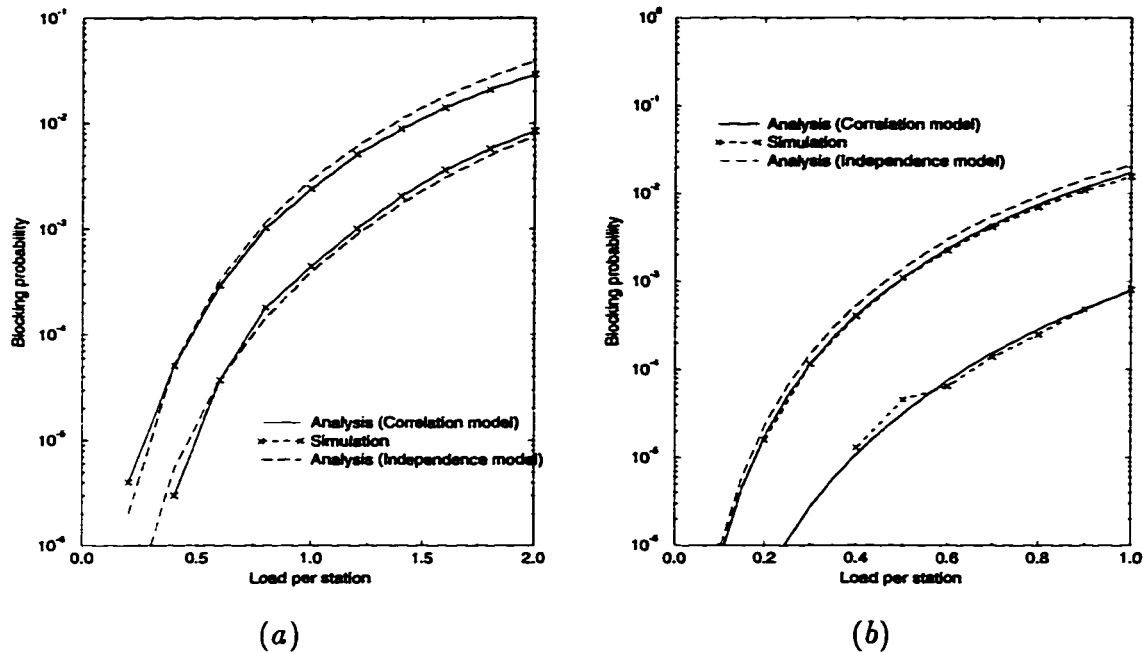


Figure 3.11: The blocking probability vs. the load per station for a hypercube network when the number of wavelengths per fiber is 5. (a) $N = 32$, and (b) $N = 1024$.

Figures 3.11(a) and 3.11(b) show how the performance varies with the load per station for a 32-node and a 1024-node hypercube, respectively. Again, the converter case curves lie below the corresponding no-converter case curves. The analysis and simulation results match very closely. Because of very low load correlation between successive links, the independence model also predicts the performance accurately, though not as well as the correlation model. Since hop-lengths are small in a hypercube network, we do not expect converters to be very useful. The figures corroborate this intuitive observation.

In Figure 3.12, we plot the blocking probability against the converter density for a 1024-node hypercube with a load of 0.1 Erlangs per station. The previous observation of converters helping more when there are more wavelengths per fiber holds for the hypercube as well. We also see that the performance improves dramatically with

an increase in the number of wavelengths. This is a direct consequence of the fact that the hypercube nodes have very high degrees, and therefore, the total number of wavelengths in the network grows very rapidly with increasing wavelengths per link.

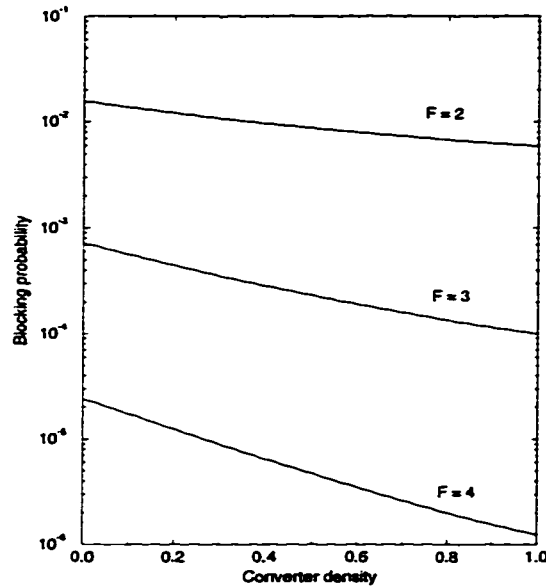


Figure 3.12: The blocking probability vs. the converter density for a 1024-node hypercube network with a load of 0.1 Erlangs per station.

The results of the last three sections lead to the remarkable observation that the usefulness of converters is a complex function of connectivity. On one hand, when connectivity is low as in the ring, hop-lengths are large and converters are expected to be more useful. On the other hand, the load correlation between successive links is high and this tends to reduce the usefulness of converters [18]. This latter effect dominates the hop-length effect in the ring. In the hypercube, the hop-lengths are small. Therefore, even though the load correlation is negligible, converters do not offer significant advantages. In the mesh network, the load correlation is fairly low while hop-lengths are large enough so that converters improve performance dramatically. The benefits of conversion are thus dependent on which of the two effects dominates in a given topology and are difficult to predict *a priori* without a detailed analysis

such as the one we have presented.

3.6 *Blocking in Random Topologies*

We studied the above networks to gain an understanding of how conversion density affects the blocking performance under varying degrees of network connectivity. In this section, we study the performance of a more likely topology for a WAN, namely an irregular topology. It is very difficult to analytically predict the performance of a given irregular topology with a particular placement of wavelength converters. However, it is possible to analyze the (ensemble) average performance of all irregular network topologies characterized by a given connectivity parameter with a converter distribution characterized by another parameter.

We employ the following model for a random topology of N nodes. Each of the possible $N(N - 1)$ directional links exists in the network with probability $\frac{b}{N-1}$, independently of all other links. Thus, each node has an average of b outgoing links. This model does not ensure that the network is connected. However, the model lends itself to easy analysis of the hop-length distribution and is used in our analysis.

We again assume that one of the shortest paths between the two end-nodes is chosen arbitrarily for a connection. The hop-length distribution is easily obtained for the regular topologies discussed but is extremely difficult to compute exactly for a random network. Exact computation of the probabilities is trivial for one and two hop paths but for longer paths, the complexity is exponential in the network size [34]. Asymptotic expressions for the shortest-path distribution are given in [35]. In our analysis, a slight variation of an approximation given in [36] is used. The approximation has been observed to be accurate for large values of N in simulations.

Given a node, let h_l be the average number of nodes that are reached on the l th hop from the given node, and let Γ_l be the average fraction of nodes (excluding the given node) that can be reached in l or fewer hops. The hop-length distribution is

computed using the following set of recurrence relations, where $p_l = h_l/(N-1)$ is the probability of an l -hop path:

$$\begin{aligned} h_l &= (N-1)(1-\Gamma_{l-1}) \left[1 - \left(1 - \frac{b}{N-1} \right)^{h_{l-1}} \right] \\ \Gamma_l &= \Gamma_{l-1} + \frac{h_l}{N-1} \end{aligned}$$

with $h_1 = b$, and $\Gamma_1 = b/(N-1)$.

The average number of exit links per node can be shown to be

$$k = \frac{(1 - \frac{1}{N-1})b}{1 - (1 - \frac{b}{N-1})^{N-2}}.$$

We assume a value of b that is sufficiently high to keep the probability of an unconnected network small. For $N \leq 500$, a value of $b \geq 10$ almost ensures connectivity.

We obtain the blocking performance using the analysis of Sections 3.1 and 3.2. Figure 3.13 shows that wavelength converters do not affect the performance much in a densely connected random network. A small increase in the number of wavelengths per link causes a significant improvement in the blocking performance, as shown in Figure 3.13. This is because of the large number of links in the network. Because of the abundance of the number of available wavelengths in the network, the blocking probability never levels off with increasing conversion density. However, the decrease in blocking with increasing conversion density is only marginal because of the short hop-lengths.

Figure 3.14 shows the dramatic improvement in performance with network connectivity. In Figure 3.15, we have plotted the blocking probability against the load per station for the no-converter and full-converter cases when the average out-degree per node is 10.

Figure 3.16 shows how the maximum number of stations that can be supported so that the blocking probability is below 10^{-3} varies with conversion density for different loads. As before, when the loads are heavy, conversion does not help much. But when

the load is light, conversion is beneficial. However, the benefits of conversion are not as great as in a less dense network such as the mesh.

The average amount of resources available per node is the average number of wavelengths available per node. For the random network considered here, this is equal to the product of the average number of links per node and the number of available wavelengths per link. Figure 3.17 shows how the blocking performance varies with the number of wavelengths per link while the network becomes less dense so that the average amount of resources in the network is held constant. We see that, when the network is fairly dense ($b \geq 10$ for a 500-node network), it is always better to have fewer links with more wavelengths per link.

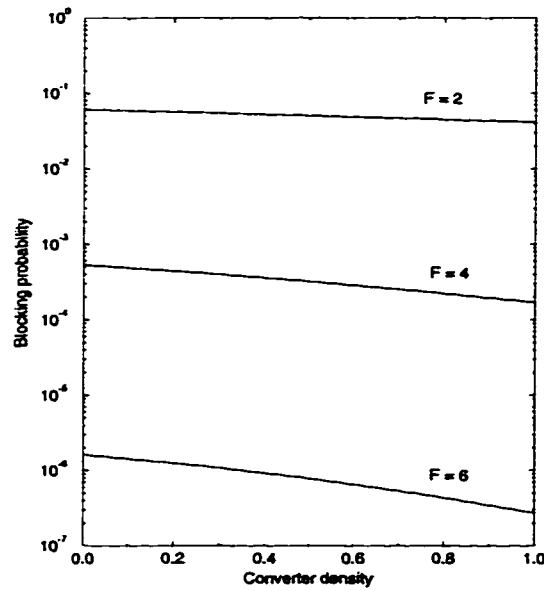


Figure 3.13: The blocking probability vs. the converter density for a random network. $N = 100, b = 10, \rho = 1$.

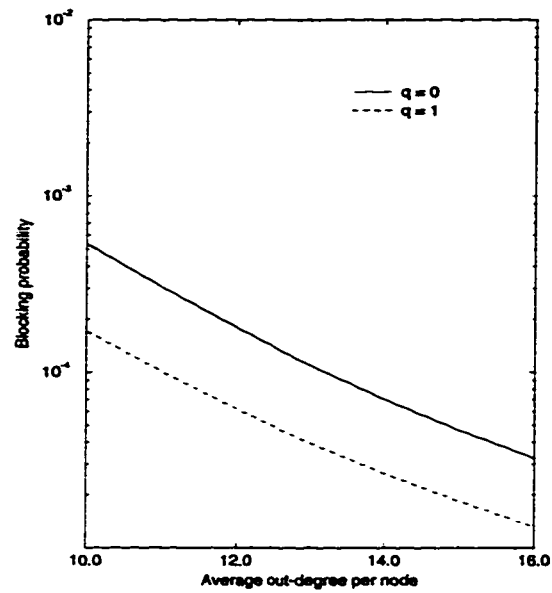


Figure 3.14: The blocking probability vs. the average out-degree per node for a random network. $N = 100$, $F = 4$, $\rho = 1$.

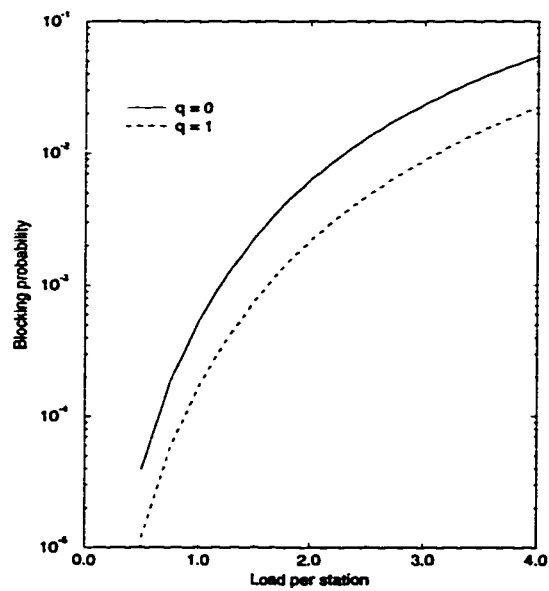


Figure 3.15: The blocking probability vs. the load per station for a random network. $N = 100$, $F = 4$, $b = 10$.

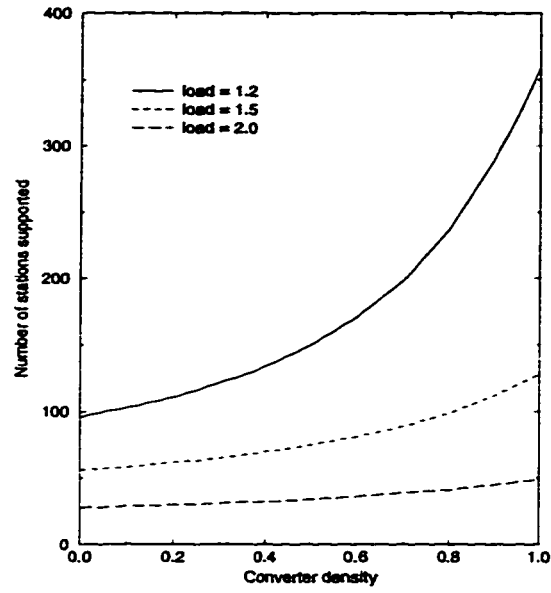


Figure 3.16: The number of stations supported vs. the converter density for a random network. $F = 4, b = 10$.

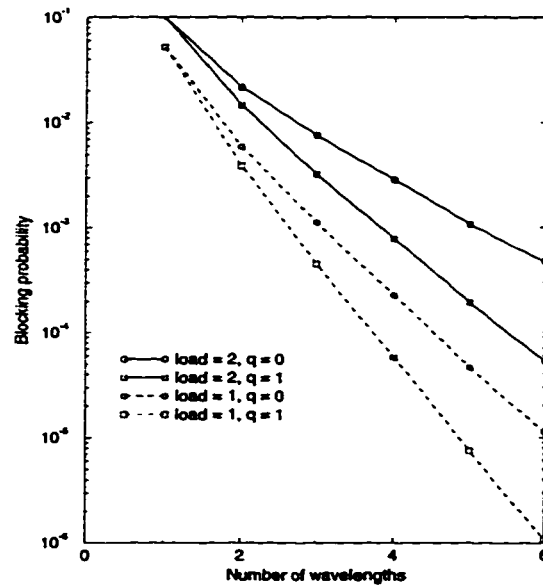


Figure 3.17: The blocking probability vs. the number of wavelengths per link for a random network. $N = 500, F \times b = 64$.

3.7 Conclusions

We have introduced the concept of sparse wavelength conversion in wavelength-routed, all-optical, circuit-switched networks. To evaluate the blocking performance of such networks, we have developed an analytical model of modest complexity. We have shown the model to be accurate for a variety of network topologies by comparing the analytical results with results from simulations. We applied this model to analyze three regular network topologies with varying degrees of connectivity – ring, mesh-torus, and hypercube.

An important conclusion of our study is that the usefulness of wavelength converters depends on the connectivity of the network in a manner that cannot be predicted by intuition. When the connectivity is low such as in the ring, converters are not very useful because of the high load correlation. This high correlation implies that the expected number of links shared by any two sessions is large; conversion merely permutes the wavelengths used by the sessions without greatly increasing the ability to accommodate a new session. When the connectivity is high such as in the hypercube or a densely connected random topology, converters are not very beneficial because of small hop-lengths, despite low load correlation and significant traffic mixing. This indicates that in a WAN with an irregular topology and a large number of links, the number of wavelengths is much more important than wavelength conversion capability. A mesh-torus network with a degree of connectivity between those of the ring and the hypercube, remarkably, offers great advantages when wavelength converters are present. This is because the link load correlation is almost negligible while the hop-lengths are large compared to the more densely connected hypercube.

The performance improves rapidly as the conversion density increases from zero, but the rate of improvement typically *decreases* with increasing conversion density. The conversion density beyond which the performance improves only marginally depends on the number of wavelengths and the network load. This suggests that placing

a converter at every node is rarely necessary, if at all. In general, wavelength converters are more effective when the number of wavelengths is larger and when the load is lower.

The results of this study point out the importance of a detailed analysis of a given network topology in determining the number and placement of wavelength converters. In the next chapter, we extend a simpler version of this model [22, 27] to the case of non-Poisson traffic and study its impact on the benefits of wavelength conversion.

Chapter 4

AN ANALYTICAL MODEL FOR SPARSE WAVELENGTH CONVERSION: NON-POISSON TRAFFIC

4.1 *Introduction*

In the previous chapter, we saw that wavelength conversion significantly improves the blocking performance in moderately dense networks such as the mesh-torus and does not have as significant an effect in sparse and very dense topologies such as the ring and the hypercube. Furthermore, we saw that typically only a small fraction of routing nodes need to be equipped with conversion capability. Those conclusions were arrived at by assuming Poisson input traffic. However, the Poisson process may not be representative of the input traffic in a wide area optical network [37].

The effect of wavelength conversion when the input traffic is non-Poisson needs to be investigated in order to obtain the sensitivity of the conclusions to the Poisson assumption. In this chapter, we study the effect of dynamic non-Poisson traffic on wavelength conversion by employing the Bernoulli-Poisson-Pascal model [29] and extending the model in [22, 27]. We evaluate the blocking probability and obtain the dependence of the blocking performance and the wavelength conversion gain on network parameters such as conversion density and peakedness of the link traffic.

In Section 4.2, we review the framework for modeling non-Poisson traffic by considering a single link. Section 4.3 presents a summary of a previous model [22, 27] that assumes link-load independence and extends it to the case of non-Poisson traffic. Finally, we use this model to analyze the blocking performance of two networks, the mesh-torus and the hypercube, in Section 4.4 and present some results. Conclusions

are presented in Section 4.5.

4.2 *The Single Link Model*

We first review the necessary background to characterize non-Poisson traffic offered to a single link.

Consider a single link with F servers and assume that a superposition of several traffic streams is offered to this link. We assume that the holding times of the arriving calls are exponentially distributed and are independent of the arrival process.

Teletraffic models with arbitrary arrival processes are $G/M/F/F$ queues, i.e., the arrival process is arbitrary, the service times are exponentially distributed, the number of servers is F , and there is no waiting room. The nature of the traffic is thus completely determined by the nature of the arrival process. Two types of arrival processes have traditionally been studied in queueing theory literature. In the first approach, the arrival process is such that the queue can be analyzed as a birth-death process where the interarrival times are exponentially distributed, but dependent on the state of the system (the number of busy servers). In the second approach, the interarrival times are independent and identically distributed. These processes, called renewal processes, are completely characterized by the distribution of the interarrival times. Note that the Poisson process is a special case of both of the above types of processes.

Traffic is often characterized by the busy circuit distribution induced by it in an infinite trunk group [29, 38]. The busy circuit distribution for a general arrival process is difficult to obtain. More importantly, the exact statistical description of the interarrival process that will be realistic in future optical networks is very difficult to ascertain. Instead of attempting to study the effect of specific non-Poisson traffic types, we employ a moment-matching technique commonly used in teletraffic theory to study overflow streams in telephone networks.

Suppose that a few moments of the offered traffic (the busy circuit distribution) are known. The moment-matching technique consists of choosing an *equivalent process* that yields the same moments and whose distribution is easily obtained, and using that equivalent process to obtain the busy circuit distribution [39, 40, 41, 42]. In practice, the number of moments that are matched vary from two to four. The two teletraffic models mentioned above are candidates for the equivalent process.

The earliest and most widely accepted moment-matching method uses two moments [39, 40] and is the Bernoulli-Poisson-Pascal (BPP) model. This method provides fairly accurate results and is attractive because of its simplicity and the ease with which the model parameters can be obtained from the first two moments. For a more detailed discussion on moment-matching techniques, the reader is referred to [29].

4.2.1 Definitions and Notation

The following definitions about the input traffic will be used in this chapter. Consider a single link with infinite servers with calls arriving at a Poisson rate λ_k^* when the system is in state k and exponentially distributed service times with mean $1/\mu$. Let p_k^* denote the probability that the infinite-server system is in state k in the steady state.

- The *offered load* is the average number of busy servers, ρ , when the traffic is offered to the link with infinite servers.
- Let σ^2 denote the variance of the number of busy servers in the infinite-server system. The *peakedness* of the input traffic, Z , is defined as the ratio of the variance to the average number of busy servers in the infinite-server system, i.e., $Z = \sigma^2/\rho$. Peakedness has been used as a measure of a traffic's deviation from Poissonian nature [40].

- A traffic is called *smooth*, *regular*, or *peaked* depending on whether Z is less than, equal to, or greater than 1, respectively. Note that a Poisson arrival process is regular.

Now, consider the link with the number of servers limited to F . Let λ_k denote the Poisson rate of arrival of calls when the system is in state k .

- The *carried load* is the average number of busy servers, ψ , when the input traffic is offered to the F -server system.
- The *unconditional occupancy distribution* is the set of probabilities, p_k , that k servers are busy in the F -server system at an arbitrary instant of time, in the steady state.
- The *arrival occupancy distribution* is the set of probabilities, p'_k , that an arriving call finds k servers busy in the F -server system, in the steady state.
- The call congestion, $B = p'_F$, is the probability that an arriving call finds the F -server system full, and therefore gets blocked.

When the arrival process is a non-Poisson process, the arrival occupancy distribution is not the same as the unconditional occupancy distribution [29].

Peakedness is an indication of traffic burstiness. In a traffic with $Z > 1$, arrivals tend to happen in clusters and tend to find the system more congested than does an observer. Therefore, in order to carry the same amount of load with peaked traffic as with Poisson traffic, the number of servers must be increased, sometimes significantly [40]. The object of this chapter is to explore the impact of traffic peakedness on the performance benefits due to wavelength conversion.

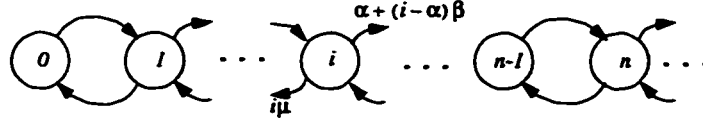


Figure 4.1: The Markov chain for the Pascal model.

4.2.2 The BPP Model

In the following discussion, we present the BPP model of [39] and compute the arrival occupancy distribution using the model.

The BPP model consists of replacing an arbitrary arrival process by another process with the same first two moments (offered load and peakedness) for which the arrival occupancy can be obtained analytically. The Bernoulli, Poisson, or Pascal model is used when Z is less than, equal to, or greater than 1, respectively, as described below.

The Pascal Model

Consider a single link with infinite servers where the interarrival times are exponentially distributed and state-dependent. Let the call arrival rate at state i , $\lambda_i^* = \alpha + (i - \alpha)\beta$ for some $\alpha > 0$ and $0 < \beta < 1$. The number of busy servers can be represented by the Markov chain shown in Figure 4.1.

The stationary probabilities of the infinite server system are given by

$$p_i^* = \left(1 - \frac{\beta}{\mu}\right)^a \left(\frac{\beta}{\mu}\right)^i \binom{a + i - 1}{i}$$

where $a \equiv \frac{\alpha}{\beta}(1 - \beta)$. This is the Pascal distribution with mean and variance

$$\begin{aligned} \rho &= \alpha \\ \sigma^2 &= \frac{\alpha}{1 - \beta}, \end{aligned} \tag{4.1}$$

where we have set $\mu = 1$. Hence, $Z = 1/(1 - \beta) > 1$, and the Pascal model can be used to model a peaked arrival process with arbitrary offered load. Given ρ and Z ,

the arrival process is modeled as the Markov process of Figure 4.1 with $\mu = 1$, $\alpha = \rho$, and $\beta = 1 - 1/Z$.

We will show later how the arrival occupancy distribution is obtained when the number of servers is restricted to F . We next proceed to the case with $Z < 1$ and discuss how the standard binomial distribution can be used as a model.

The Bernoulli Model

Let a single link with infinite servers be offered traffic from a finite number of sources n . Consider the Markov chain in Figure 4.2 for some $\lambda, \mu > 0$.

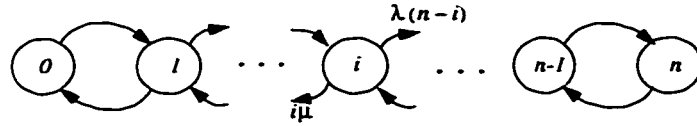


Figure 4.2: The Markov chain for the Bernoulli model.

The stationary probabilities are given by

$$p_i^* = \binom{n}{i} \left(\frac{\lambda}{\lambda + \mu} \right)^i \left(\frac{\mu}{\lambda + \mu} \right)^{n-i}$$

which is the Bernoulli (the binomial) distribution¹. The mean and peakedness of the offered traffic are given by

$$\begin{aligned} \rho &= n \frac{\lambda}{\lambda + \mu} \\ Z &= \frac{\mu}{\lambda + \mu}. \end{aligned} \tag{4.2}$$

Hence, this can be used as a model for smooth traffic. Given an offered load ρ and a peakedness $Z < 1$, the arrival process is assumed to be the process of Figure 4.2

¹ Note that in [18], the arrival occupancy distribution in the finite-server system was modeled as a binomial. Here, we are using a binomial to model the unconditional distribution in an infinite-server system.

with $\mu = 1$, and

$$\begin{aligned}\lambda &= \frac{1-Z}{Z} \\ n &= \frac{\rho}{1-Z}.\end{aligned}\tag{4.3}$$

Note that when n above is not an integer, the physical interpretation of n as the number of sources fails to hold; nevertheless, the model has been found to yield sufficiently accurate numerical results for congestion probabilities in telephony [39].

A limitation of the Bernoulli model when the number of servers is finite is that the number of sources n must be at least equal to the number of servers F . This places a restriction on the range of Z ; in particular, $Z \geq 1 - \rho/F$.

Having presented models to approximate smooth, regular, and peaked traffic, we now proceed to derive the arrival occupancy distribution when these traffics are presented to a link where the number of servers is limited to F . This distribution will be used in analyzing the blocking performance in Section 4.3.

Arrival Occupancy Distribution

Suppose an arrival process is offered to a single link with F servers. We derive below the arrival occupancy distribution for a general arrival process with state-dependent exponential interarrival times. Let λ_i denote the arrival rate of calls when the system is in state i . Let the departure rate in state i be μ_i . The BPP model further assumes [39] that the unconditional occupancy distribution in the F -server system is a truncated Pascal, Poisson, or binomial distribution. In order to achieve this, we assume that $\lambda_i = \lambda_i^*$, $i = 0, 1, \dots, F-1$.

Let P_{ji} denote the conditional probability that an arrival in steady state finds the system in state i given that the previous arrival found the system in state j . The arrival occupancy distribution is the solution to the system of linear equations

$$p'_i = \sum_{j=0}^F p'_j P_{ji}, \quad i = 0, 1, \dots, F.\tag{4.4}$$

We now determine P_{ji} . Given that the system was in state j when the previous call arrived, the system is in state $j + 1$ immediately after that arrival (except when the system was in state F in which case the state is unchanged). The probability that the next call finds the system in state i is exactly the probability that $j + 1 - i$ calls depart the system before the new arrival occurs.

Using the properties of the exponential distribution, we can write

$$P_{ji} = \begin{cases} \frac{\mu_{j+1}}{\lambda_{j+1} + \mu_{j+1}} P_{j-1,i}, & 0 \leq i \leq j \leq F-1, \\ \frac{\lambda_i}{\lambda_i + \mu_i}, & 1 \leq i \leq F-1, j = i-1, \\ \frac{\mu_{i+1}}{\lambda_{i+1}} \frac{\lambda_i}{\lambda_i + \mu_i} P_{F,i+1}, & 0 \leq i \leq F-1, j = F, \\ \frac{\lambda_F}{\lambda_F + \mu_F}, & i = F, j = F, F-1, \\ 0, & \text{otherwise,} \end{cases} \quad (4.5)$$

with the understanding that $P_{-1,0} = 1$.

After some manipulations involving (4.4) and (4.5), we obtain

$$p'_i = \frac{\lambda_i}{\mu_i} p'_{i-1}, \quad i = 1, 2, \dots, F. \quad (4.6)$$

p'_0 can be determined from $\sum_{i=0}^F p'_i = 1$. Note the difference between the arrival occupancy distribution and the unconditional occupancy distribution which satisfies

$$p_i = \frac{\lambda_{i-1}}{\mu_i} p_{i-1}, \quad i = 1, 2, \dots, F. \quad (4.7)$$

It is seen that for a Poisson arrival process the two distributions are the same.

It is important to observe that the unconditional occupancy probabilities are independent of λ_F whereas the arrival occupancy probabilities are not. When the infinite-server system is truncated to an F -server system, the carried load ψ can be computed using the unconditional occupancy distribution as $\psi = \sum_{i=0}^F i p_i$. Moreover, the carried load ψ and offered load ρ are related by

$$\psi = \rho(1 - B) \quad (4.8)$$

where B is the probability that an arriving call finds all F servers busy. Since $B = p'_F$, (4.6) and (4.8) define a relation between λ_F and the unconditional occupancy distribution. It can be shown using $\rho = \sum_{i=0}^{\infty} i p_i^*$ that $\lambda_F = (\sum_{i=0}^{\infty} p_i^* \lambda_i^* - \sum_{i=0}^{F-1} p_i \lambda_i) / p_F$, where λ_i^* is the arrival rate at state i in the infinite-server system.

4.3 Network Blocking Analysis

In this section, we summarize the network model of [22, 27] and extend it to the case of non-Poisson traffic using the model presented in Section 4.2.

4.3.1 Model assumptions

The following assumptions, similar to those in the previous chapter, are used in our model.

1. Calls arrive at each node according to a random process and the offered load is ρ_s . Each call is equally likely to be destined to any of the remaining nodes.
2. Call holding times are exponentially distributed with unit mean.
3. A shortest-path is chosen randomly to route a call.
4. Each call requires a full wavelength on each link it traverses.
5. The number of wavelengths, F , is the same on all links. Each node is capable of transmitting and receiving on any of the F wavelengths.
6. A call is blocked if the chosen path is blocked.
7. Wavelengths are assigned to a session randomly from the set of free wavelengths on the associated path.

8. Link loads are statistically independent².

4.3.2 Networks without Conversion

Throughout this section, all events regarding the status of wavelengths refer to events occurring when a call arrives. A wavelength is defined as “free” on a path if that wavelength is not used on any of the links constituting the path. A wavelength is “busy” on the path otherwise. First, consider the case of networks with no wavelength conversion. Given the load per station ρ_s , the load per link ρ can be computed using the uniform traffic and shortest-path routing assumptions as

$$\rho = N\rho_s\bar{H}/L, \quad (4.9)$$

where N is the number of nodes, $\bar{H}$ is the average path length and L is the number of links in the network. The arrival occupancy distribution on a single link can be computed using this value of ρ for various values of the link peakedness Z using the analysis of Section 4.2. Note that the actual load carried by each link will be less than ρ due to blocking (which in turn depends on ρ). As we observed in Chapter 3, ignoring this feedback effect does not produce significantly different results when the blocking probabilities are low (on the order of 10^{-3}), while considerably decreasing the computation time.

To compute the blocking probabilities in networks without conversion, the distribution of free wavelengths is required. In a two-hop path, the probability that n wavelengths are free on the path is obtained as [22, 27]

$$t_n^{(2)} = \sum_{i=0}^F \sum_{j=0}^F R(n|i, j) p'_{F-i} p'_{F-j}, \quad (4.10)$$

² This assumption has been shown to yield accurate blocking probability estimates for Poisson traffic for the topologies we study in this chapter [28].

where $R(n|i, j)$ is defined as

$$R(n|i, j) = \frac{\binom{i}{n} \binom{F-i}{j-n}}{\binom{F}{j}}, \quad (4.11)$$

for $\max(0, i + j - F) \leq n \leq \min(i, j)$ and is 0 otherwise.

The free wavelength distribution on an l -hop path can be obtained recursively using the distribution for the two-hop path by considering the first $l - 1$ hops as the first hop and hop l as the second hop of a two-hop path. The probability of finding n free wavelengths on an l -hop path is then written as

$$t_n^{(l)} = \sum_{i=0}^F \sum_{j=0}^F R(n|i, j) t_i^{(l-1)} p'_{F-j}, \quad (4.12)$$

where p'_i is as obtained in Section 4.2. The probability of blocking on an l -hop path is given by $t_0^{(l)}$.

4.3.3 Networks with Conversion

We model networks with wavelength conversion by assuming that each node is capable of converting any wavelength to any other with probability q , independently of the other nodes, as in the previous chapter. This leads to a binomial distribution with mean Nq on the number of converter nodes in the network. The blocking performance we obtain is the ensemble average of the blocking probability over this distribution.

Consider an l -hop path and let the nodes on the path be numbered $s, 1, 2, \dots, l - 1, d$ as shown in Figure 3.3. Let $P_b^{(i)}(q)$ denote the probability of blocking on an i -hop path. $P_b^{(1)}(q)$ is clearly p'_F . $P_b^{(l)}(q)$ can be computed recursively in terms of the blocking probabilities on paths of fewer hops and on paths with no converters as

$$P_b^{(l)}(q) = t_0^{(l)}(1 - q)^{l-1} + \sum_{i=1}^{l-1} [1 - (1 - P_b^{(i)}(q))(1 - t_0^{(l-i)})] q(1 - q)^{l-i-1}. \quad (4.13)$$

The first term is the probability that there are no converters in the intermediate nodes along the given l -hop path times the probability of blocking in this case (which

is $t_0^{(l)}$). The i th term in the series is the probability that node i is the last converter in the given l -hop path times the probability of blocking in this case. The probability that a call is not blocked in this case is the product of the probability that it is not blocked on the i -hop path (which is $1 - P_b^{(i)}(q)$) and the probability that it is not blocked on the last $l - i$ hops given there are no wavelength converters on the last $l - i$ hops (which is $1 - t_0^{(l-i)}$). The network blocking probability is then computed as

$$P_b(q) = \sum_{l=1}^{N-1} P_b^{(l)}(q) r_l \quad (4.14)$$

where r_l is the probability that an l -hop path is chosen to route an arriving call.

Given a network topology with an offered load per station of ρ_s , we first determine the hop-length distribution r_l . The blocking performance for different values of link traffic peakedness can then be studied by using the analysis presented in Section 4.2 and this section. In the next section, we evaluate the effect of link peakedness on call blocking performance and wavelength conversion for two regular network topologies that have different degrees of connectivity – the mesh-torus and the hypercube.

4.4 Performance Results

We study the effects of traffic peakedness on networks with sparse wavelength conversion in two different ways. In the first approach, the blocking probability of a given network is evaluated as a function of the peakedness and conversion density. In the second approach, we define the *conversion gain* G_q as the ratio of the offered load that can be supported with a given value of conversion density q to the load that can be supported without conversion, for a given blocking probability. The evaluation is performed on two network topologies – an 11×11 mesh-torus and a 128-node binary hypercube, both with 10 wavelengths per fiber. The mesh-torus has four outgoing fibers per node and the hypercube has $7 (= \log_2 128)$ outgoing fibers per node. The routing is shortest-path routing where each of the possible paths is chosen with equal probability.

First, for the mesh network, we plot the blocking probability against the offered load per station ρ_s for link peakedness values of 1.0, 1.25, and 1.5 in Figure 4.3. Two observations can be made from these figures. One is that for a given value of peakedness, the curves for no conversion and full conversion are essentially parallel indicating that the full conversion gain is not a sensitive function of the blocking probability. More importantly, as the peakedness increases, the supportable loads for a fixed blocking probability rapidly decrease.

Note that while the link load is simply given by the sum of the node traffic loads that are superposed on the link, the link peakedness depends on the exact proportions of the node traffic that are superposed on the link. This results in an interesting relationship between the peakedness of link traffic and that of node traffic.

Let ρ_s and σ_s^2 be the mean and variance of the offered traffic at each station. Consider an arbitrary link l and let s_{il} be the fraction of the traffic originating at node i that flows through l . Then, the mean ($\rho(l)$) and variance ($\sigma^2(l)$) of the traffic offered to link l can be computed using simple properties of the expectation of a random variable as

$$\rho(l) = \sum_{i=1}^N s_{il} \rho_s \quad (4.15)$$

$$\sigma^2(l) = \sum_{i=1}^N [s_{il}(1 - s_{il})\rho_s + (1 - s_{il})^2\sigma_s^2]. \quad (4.16)$$

Using (4.16), the peakedness of traffic offered to link l , $Z(l)$, is given in terms of the peakedness of node traffic Z_s as

$$Z(l) = 1 + A(l)(Z_s - 1),$$

where

$$A(l) \stackrel{\text{def}}{=} \frac{\sum_{i=1}^N s_{il}^2}{\sum_{i=1}^N s_{il}}.$$

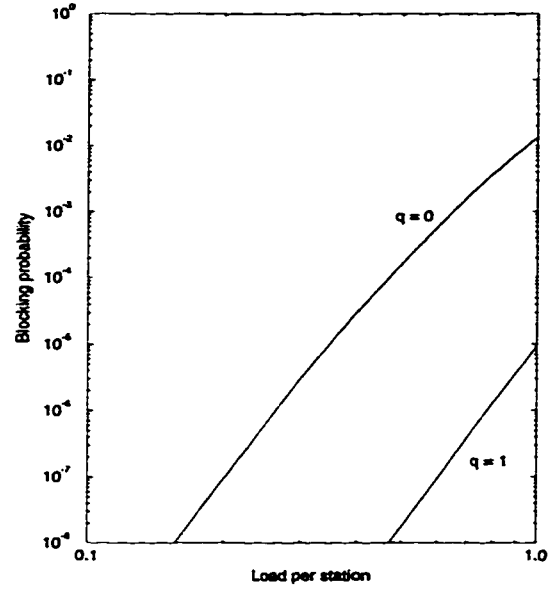
Given a topology and routing algorithm, the fractions $\{s_{il}\}$ have to be determined to compute $Z(l)$ and $\rho(l)$ for each link l . Note from the definition of $A(l)$ that it is

no greater than 1, and therefore, link peakedness is never more than peakedness of traffic offered to a node. Also note that $A(l)$ is smaller when there is more traffic mixing.

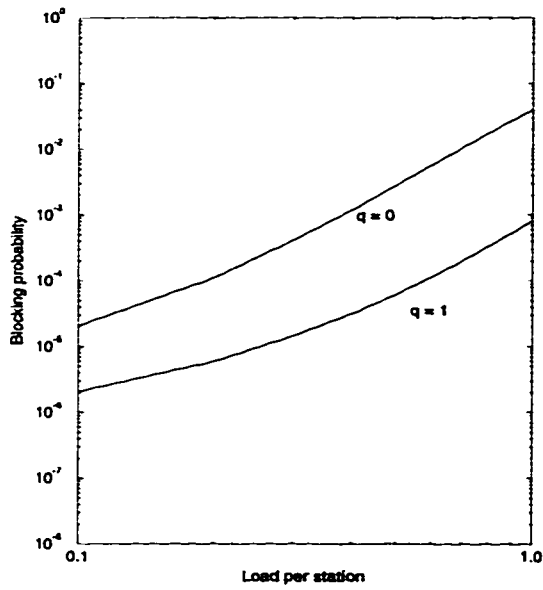
When the topology is a mesh-torus and the routing is random shortest-path routing, the fractions $\{s_{il}\}$ are the same for all links, and can be determined by counting the number of shortest paths between two nodes that use a specific link. For a 11×11 mesh-torus, $A(l)$ turns out to be 0.080347 for all links l . Thus, the station peakedness values that correspond to link peakedness values of 1, 1.25, and 1.5 are, respectively, 1, 4.11, and 7.22.

Next, we show the blocking probability as a function of the conversion density for the mesh and the hypercube in Figures 4.4(a) and 4.4(b), respectively. As pointed out earlier by several researchers [18, 22, 27, 28], the blocking probability decreases significantly with Poisson input traffic ($Z = 1$) when wavelength converters are employed. From the figures, we observe that as the traffic becomes more peaked, the blocking probabilities do not diminish as significantly as for Poisson traffic in either topology. This suggests that peaked traffic decreases the performance gain with wavelength conversion if it is measured as the decrease in blocking probability with wavelength conversion.

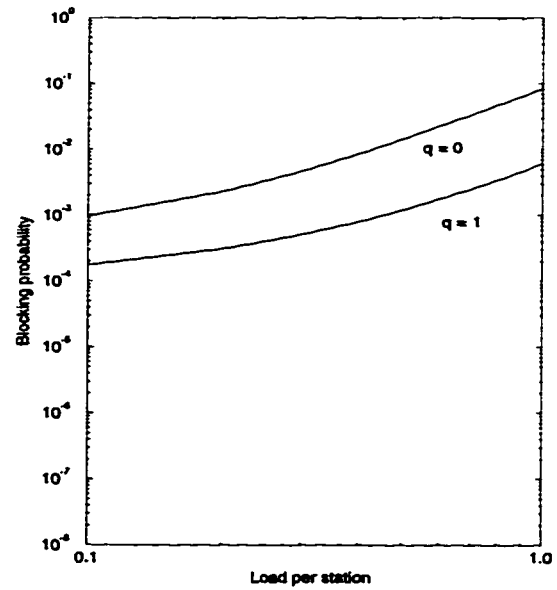
The graphs can also be used to determine the amount of conversion needed for a given traffic peakedness and blocking probability. For example, in the mesh with Poisson input traffic, a conversion density of approximately 40% is sufficient for realizing a blocking probability of 10^{-3} whereas when the link peakedness increases to 1.25, close to 100% conversion is required. The blocking probability is plotted as a function of link peakedness for different values of conversion density in the hypercube with a load of $\rho_s = 1.0$ per station in Figure 4.5. It is observed that the probability of blocking increases dramatically as the link traffic becomes more peaked. The situation is similar in the mesh (not shown). This is because of the fact that an arrival sees a much more congested system with peaked traffic due to its “bursty” nature.



(a)

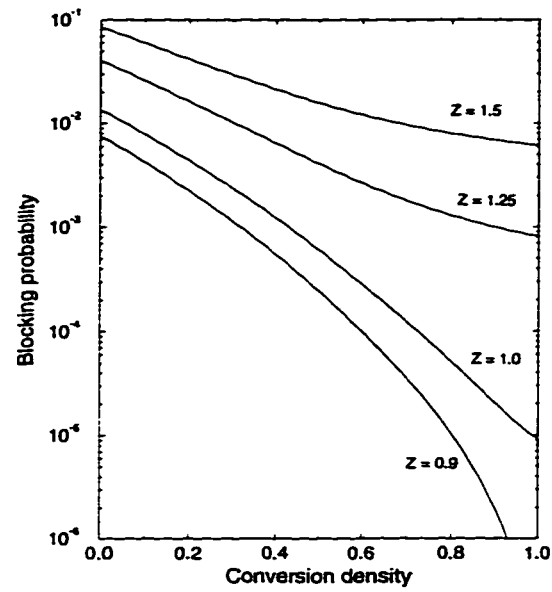


(b)

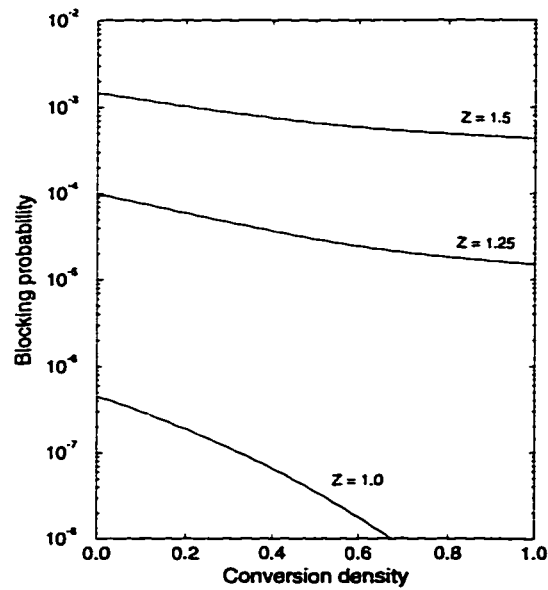


(c)

Figure 4.3: The blocking probability vs. the load per station for the mesh network with no conversion and full conversion. (a) $Z = 1.0$ (b) $Z = 1.25$ (c) $Z = 1.5$.



(a)



(b)

Figure 4.4: The blocking probability vs. the conversion density with a load of 1.0 per station for (a) the mesh, and (b) the hypercube.

Finally, the conversion gain is plotted against link peakedness in Figure 4.6 for different values of conversion density in the mesh network. The blocking probability is maintained at 10^{-3} . The graphs show that the gain is fairly insensitive to link peakedness up to a certain value (about 1.45 in this case) for any value of conversion density. Beyond this point, the gain dramatically increases. However in the region of large Z , the loads that can be supported with and without conversion are small (e.g., for $Z = 1.55$, the supportable station loads with no conversion and full conversion are 0.015 and 0.297, respectively). Therefore even a modest increase in supportable load with conversion translates to a huge conversion gain (which is a ratio). For the hypercube (not shown here) the breakaway point occurs at a higher peakedness.

4.5 Conclusions

In this chapter, we have developed a methodology for studying the benefits of wavelength converters for non-Poisson dynamic traffic. We used the Pascal model which has been used by teletraffic engineers to characterize the non-Poisson nature of overflow traffic from a trunk. We derived the arrival occupancy distribution for this model and utilized it to analyze the effect of traffic peakedness on wavelength conversion benefits and call blocking performance.

For the two topologies considered, our models predict that for a fixed load, peakedness dramatically increases the blocking probability, and wavelength conversion can only mildly alleviate performance degradation. If the metric of interest is the reduction in blocking probability with wavelength conversion, then traffic peakedness reduces the benefits of wavelength conversion (Figure 4.4). However, if the performance metric is the increase in network throughput achievable with wavelength conversion at a fixed blocking probability, peakedness does not affect the benefits of wavelength conversion for a large range of peakedness values (Figure 4.6). In this case, wavelength conversion assumes a more significant role at large peakedness values.

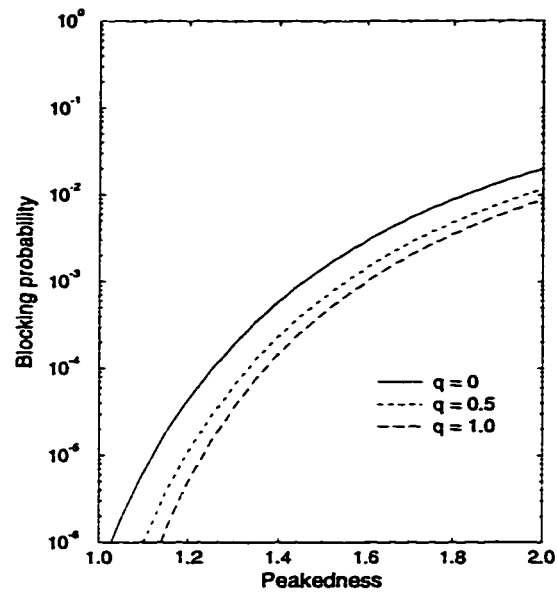


Figure 4.5: The blocking probability vs. the link peakedness for the hypercube. $\rho_s = 1$.

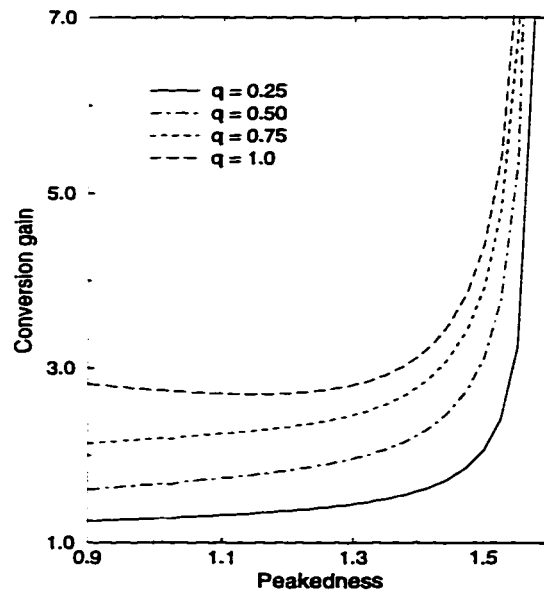


Figure 4.6: The conversion gain vs. the link peakedness for the mesh network. $P_b = 10^{-3}$.

Chapter 5

ON THE OPTIMAL PLACEMENT OF WAVELENGTH CONVERTERS

5.1 Introduction

In Chapters 3 and 4, we evaluated the ensemble average blocking performance of networks with a random number of wavelength converters randomly placed in the network. In this chapter, we consider a closely related and more practical problem. Given a network topology, a certain number of converters, and traffic statistics between the nodes, the problem of interest is the optimal placement of converters in the network. To our knowledge, there does not exist any published work that addresses this problem. There are several approaches to this problem depending on the optimality criterion and the traffic model.

In one approach, the maximum offered traffic (the number of sessions) λ_{sd} for all node pairs (s, d) is given. Let M_{sd} be the actual number of sessions established between s and d . The problem is to determine the converter locations so that $\sum_{s,d} M_{sd}$ is maximized subject to the wavelength continuity constraint on segments between successive converters and the traffic demand constraint, i.e., $M_{sd} \leq \lambda_{sd}$. Note that this problem contains the routing and wavelength assignment (RWA) problem as a subproblem. The placement problem can be posed as an integer-linear programming problem (see Appendix A) similar to the formulation given in [10], which in fact is a special case of the problem at hand. Previous work has shown that this special case is an NP-complete problem for arbitrary topologies, and even for rings [24]. Note, however, that in the case of a bus network, the converter placement problem is trivial

in this formulation. This is because any set of sessions that can be accommodated with full wavelength conversion (where every node has a wavelength converter) can also be accommodated with no wavelength conversion. This follows from the properties of an interval graph [43].

In this chapter, we assume a dynamic traffic model in which connections arrive and depart from the network in a random manner and the goal is to place the given number of converters in the network such that the network call blocking probability is minimized. We present a dynamic programming solution for minimizing the blocking probability on a path. This solution is then generalized to bus and ring networks. Our solution to the problem can be used in conjunction with any performance model for dynamic traffic. For arbitrary topologies, heuristics based on the approaches presented here may need to be used.

There are several factors which affect the optimal solution to the converter placement problem. Intuition would suggest that converters be placed at nodes which process the highest amount of transit traffic. However, placing a converter at a node that has a high transit traffic rate but does very little mixing (or switching) of traffic may not be the optimal strategy. On the other hand, if the transit traffic rate at a node is very low, then the optimal strategy may not place a converter at that node even if it mixes a significant amount of traffic. Furthermore, the distances between converters are likely to affect the optimal placement. As the distance between converters increases, the blocking probability increases. Since the number of available converters is limited, a judicious placement of converters that balances these and other effects is necessary.

The rest of this chapter is organized as follows. In Section 5.2, we consider a single path of length H and assume link load independence. The motivation for studying the placement problem on a path is the following. Converter placement in an arbitrary network is a hard problem in general and sub-optimal heuristics would probably be required. In such a case, we envisage the placement process to proceed

in multiple phases where at each phase a certain number of converters are placed optimally on a path. Note, however, that because of the interaction between the paths, the phases may need to be iterated over to obtain a satisfactory solution. It may also be necessary to place converters on longer paths in the network to provide quality of service guarantees.

In Section 5.2, we first consider uniform link loads on the path and show that uniformly spaced converters minimize the blocking probability of an end-to-end call. We also derive a recurrence relation for the blocking performance when the converters are placed randomly along the path. We then obtain an approximate expression for the utilization gain achievable with uniform/optimal placement relative to random placement and show that significant gains are achievable with a good placement policy. Next, we relax the uniform link load assumption and obtain the optimal solution via dynamic programming. In the last part of the section, we extend the optimality criterion to take into account the average blocking probability (P_b) over all source-destination $((s, d))$ pairs along the path.

In Section 5.3, we consider a bus network. The difference between this and the path considered in Section 5.2 is that the traffic on the links is entirely due to the traffic originating at the nodes of the bus network, whereas the link traffic in the path is possibly due to merging traffic from other parts of the network. The low connectivity of the bus network induces a strong correlation among link loads. This is taken into account by modifying the performance model accordingly. The solution in Section 5.3 can also be used to optimally place converters in a path of an arbitrary network where link load independence is not justified. In Section 5.4, we extend the dynamic programming solution for the path to obtain the optimal converter placement for the ring topology.

The blocking performance model we use primarily in this chapter is the wavelength independence model of [18] which we refer to as the *binomial model* because

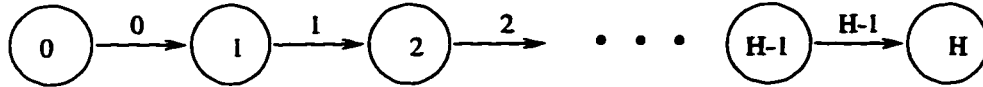


Figure 5.1: A path of length H .

the number of busy wavelengths on a link is assumed to be binomially distributed¹. However, as mentioned earlier, the solution technique is independent of the performance model and we compare some of the results obtained with the binomial model and the Poisson model of Chapter 3 in Section 5.5. Conclusions are presented in Section 5.6.

5.2 Converter Placement in a Path

In this section, we consider a path of length H with negligible correlation between link loads. We first assume uniform link loads and derive an approximate expression for the utilization gain achievable by optimal converter placement relative to a random placement of a given number of converters. We start by considering an end-to-end call on a path of length H shown in Figure 5.1. Let the nodes along this path be numbered $0, 1, \dots, H$, and let the link loads per wavelength be ρ_i , $i = 0, \dots, H - 1$. That is, the probability that a given wavelength is occupied on link i is ρ_i , and the wavelength occupancy events are assumed to be statistically independent of other wavelengths on the same link and on other links. (This is the binomial model of [18].)

¹ Note that this model is different from the binomial model of Chapter 4 in which it is the steady-state distribution of busy wavelengths on a link that is binomial as opposed to occupancy distribution upon arrival.

5.2.1 Uniform Link Loads

First suppose that $\rho_i = \rho$, $i = 0, 1, \dots, H - 1$. The number of wavelengths on each link is assumed to be F and each link has a single fiber. The performance model we use is obtained by assuming that each wavelength is used on a link with probability ρ independently of the other wavelengths [18]². The goal is to place K converters among the $H - 1$ intermediate nodes of the path so that the blocking probability of an end-to-end call (a call from node 0 to node H) is minimized.

Define a *segment* to be the set of links between two consecutive converter nodes. Let l_i be the hop-length of the i th segment, $i = 0, \dots, K$ (nodes 0 and H can be assumed to contain dummy converters). Let $\mathbf{L}_K = (l_0, l_1, \dots, l_{K-1}, l_K)$, where $\sum_{i=0}^K l_i = H$, be the vector denoting the hop-lengths of the $K + 1$ segments, called the *length vector* henceforth. The success probability of the call in a subsegment³ of length l , denoted by $f(l)$, is given by

$$f(l) = 1 - (1 - \bar{\rho}^l)^F \quad (5.1)$$

where $\bar{\rho} = 1 - \rho$.

The success probability of the end-to-end call when the length vector is $\mathbf{L}_K$ is denoted by $S(\mathbf{L}_K)$, and is given by

$$S(\mathbf{L}_K) = \prod_{i=0}^K f(l_i).$$

Let the length vector in the optimal placement be $\mathbf{L}_K^{\text{opt}}$, i.e.,

$$S_{\text{opt}}(H, K) = S(\mathbf{L}_K^{\text{opt}}) = \max_{\mathbf{L}_K} S(\mathbf{L}_K).$$

² The average number of busy wavelengths on a link using this model is $F\rho$, the same as would be obtained by assuming Poisson traffic and ignoring the effects of blocking on carried load. This model will serve our purpose here. More accurate numerical results for Poisson traffic could be obtained by choosing the wavelength occupation probability as in [20]. The assumption of random wavelength assignment to an arriving connection is implicit in this model.

³ Any set of consecutive links of a segment is called a subsegment.

where $S_{\text{opt}}(H, K)$ denotes the success probability of an end-to-end call when K converters are optimally placed along an H -hop path. Suppose $(K + 1)$ divides H . We show below that $l_i^{\text{opt}} = H/(K + 1)$, $i = 0, 1, \dots, K$, is the optimal length vector. To prove this, we need to show that $\prod_{i=0}^K f(l_i^{\text{opt}}) = [f(\frac{H}{K+1})]^{K+1} \geq S(\mathbf{L}_K)$ for any feasible $\mathbf{L}_K$. Equivalently, we must show that $\ln f(\frac{H}{K+1}) \geq \frac{1}{K+1} \sum_{i=0}^K \ln f(l_i)$ where the l_i 's are the elements of a feasible length vector. This follows from the lemma below which shows that $\ln f(l)$ is a concave function of the continuous variable $l \in (0, H)$. We remark here that the concavity of $\ln f(l)$ is sufficient but not necessary for the optimality of uniform placement.

Lemma: $\ln f(l)$ is a concave function of $l \in (0, H)$.

Proof: Let $g(l) = \ln f(l)$. We have the second derivative

$$g''(l) = \frac{f''(l)f(l) - [f'(l)]^2}{f^2(l)}. \quad (5.2)$$

Evaluating $f(l)$, $f'(l)$, and $f''(l)$ and substituting in (5.2), we obtain

$$g''(l) = \frac{F(1 - \bar{\rho}^l)^{F-2}(\ln \bar{\rho})^2(1 - F\bar{\rho}^l - (1 - \bar{\rho}^l)^F)}{f^2(l)}.$$

Now notice that for every $0 \leq \rho \leq 1$ and $F \geq 1$, $(1 - \bar{\rho}^l)^F + F\bar{\rho}^l$ is a non-decreasing function of l and is lower bounded by 1. Since the remaining factors are all positive, $g''(l) \leq 0$. ■

The success probability of an end-to-end call under optimal placement is therefore

$$\begin{aligned} S_{\text{opt}}(H, K) &= \left[f\left(\frac{H}{K+1}\right) \right]^{K+1} \\ &= \left[1 - \left(1 - \bar{\rho}^{\frac{H}{K+1}}\right)^F \right]^{K+1} \end{aligned} \quad (5.3)$$

It must be noted that this expression is exact only if $(K + 1)|H$. When this is not the case, the integral constraint on the segment lengths will make the expression in (5.3) an upper bound on the actual success probability. In that case, the optimal strategy is to place the converters as uniformly as possible, i.e., converters are placed so that there are $y = H - a(K + 1)$ segments of length $a + 1$ and $K + 1 - y$ segments of length a where $a = \lfloor \frac{H}{K+1} \rfloor$.

Random Placement

We now derive a recurrence relation for the success probability of the end-to-end call when the K converters are placed randomly on the path, i.e., each of the $\binom{H-1}{K}$ converter placement configurations is equally likely. Let us define $S_r(m, j)$ to be the success probability of an end-to-end call when j converters are randomly placed along a path of length m . Then, by conditioning on the position of the converter that is placed at the lowest node index, we obtain the recursive relationship

$$S_r(m, j) = \sum_{i=1}^{m-j} f(i) S_r(m-i, j-1) \frac{\binom{m-i-1}{j-1}}{\binom{m-1}{j}} \quad (5.4)$$

for $1 \leq j \leq m-1$. Obviously $S_r(m, 0) = f(m)$. In (5.4), $\frac{\binom{m-i-1}{j-1}}{\binom{m-1}{j}}$ is the probability that the converter with lowest node index is placed at node i . In obtaining (5.4), we used the fact that all $\binom{m-i}{j-1}$ possible converter placement configurations, given that node i is the lowest index node with a converter, are equiprobable.

An approximation to $S_r(H, K)$ can be obtained by assuming that each of the $H-1$ intermediate nodes is a converter with probability $q = \frac{K}{H-1}$ independently of the other nodes, as we did in Chapters 3 and 4. As before, we will refer to the parameter q as the converter (or conversion) density of the network. Let this approximation be denoted by $S_a(H, q)$. Then, a recursive relation for $S_a(m, q)$ is

$$S_a(m, q) = (1-q)^{m-1} f(m) + \sum_{i=1}^{m-1} f(i) S_a(m-i, q) q (1-q)^{i-1} \quad (5.5)$$

for $0 < q \leq 1$, and $S_a(m, 0) = f(m)$. $S_a(H, q)$ gives the ensemble average blocking probability over a binomially distributed number of converters, with expected value $K = q(H-1)$, randomly placed among the $H-1$ nodes. We will show that the approximation is reasonable by comparing $S_r(H, K)$ and $S_a(H, q)$ later. Note that $S_a(m, q)$ is easier to compute than $S_r(m, j)$ given in (5.4), because (5.4) is a recursion in the two variables m and j , while (5.5) is a recursion in one variable (m) only. We will use this approximation to obtain an approximate expression for the utilization gain, to be defined shortly.

To see the benefits of optimal converter placement with respect to random placement, we plot the end-to-end blocking probability (P_b) as a function of ρ for $H = 10$ and $F = 10$ and $K = 1, 2$ in Figure 5.2. (Unless otherwise specified, all numerical results presented in this chapter are for $H = 10$ and $F = 10$.) We observe that the performance with two converters randomly placed is poorer than the performance with one converter optimally placed. Also notice that the performance difference between optimal and random placements is slightly enhanced as K increases from 1 to 2.

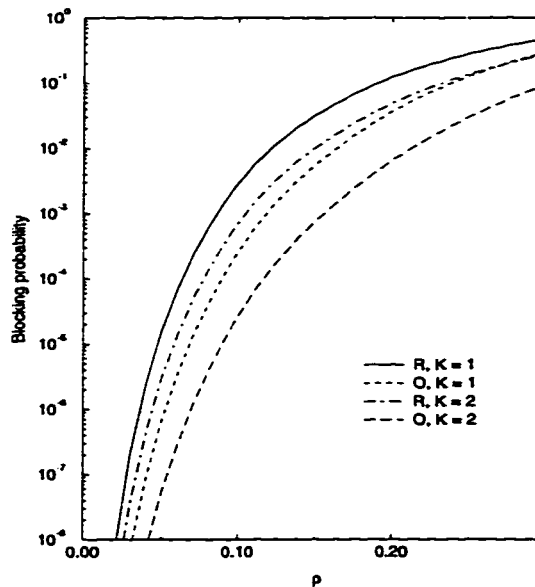


Figure 5.2: P_b vs. ρ for optimal/uniform (O) and random (R) placement with $K = 1$ and 2 converters.

The performance effect of the number of converters can be seen more clearly in Figure 5.3 where P_b is plotted against K . Notice that conversion plays a significant role in improving the blocking probability. This is of course due to link load independence and the long path length, as predicted by previous models [18, 22] and our model of Chapter 3. More importantly, it can be seen that there is a large improvement in performance between random and optimal placement, reaching a maximum

of more than 2 orders of magnitude at $K = 4$. Also plotted in Figure 5.3 is the performance when a random number of converters (with mean K) are placed randomly on the path. This is the model we assumed in the last two chapters because of the difficulty in obtaining the performance with an exact number of converters placed in a network of arbitrary topology. We observe that this performance approximates the performance of random placement of a fixed number of converters reasonably well. We will find that the ensemble average P_b estimate is consistently higher than the P_b with an exact number of converters randomly placed. As expected, all curves converge at $K = 0$ and $K = H - 1$. Notice that the R_r and O curves also converge at $K = H - 2$. This is because random placement of $H - 2$ converters produces K segments of length 1 and one segment of length 2. Since link loads are uniform, the location of this 2-hop segment does not matter and the performance is exactly the same as obtained by optimally placing $H - 2$ converters. The effect of the length of the maximum segment on P_b is worth mentioning here. Consider the O curve. The length of the maximum segment under optimal (uniform) placement as K increases from 0 to 9 are, respectively, 10, 5, 4, 3, 2, 2, 2, 2, 2, 1. It can be seen from Figure 5.3 that there is a significant performance change when the maximum segment length changes, and the performance improvement when the maximum segment length does not change (from $K = 4$ to 8) is only marginal.

Another metric to study the importance of converter placement is the utilization gain, G_u^r [18]. Let the target blocking probability P_b , and the number of converters K be given. Let ρ_o be the maximum load per link per wavelength achievable for the given P_b with the converters optimally placed, and let ρ_r be the maximum load achievable with the converters randomly placed. Then the utilization gain is defined as

$$G_u^r \stackrel{\text{def}}{=} \frac{\rho_o}{\rho_r}.$$

Also, let ρ_q be the maximum load when a random number of converters (with an

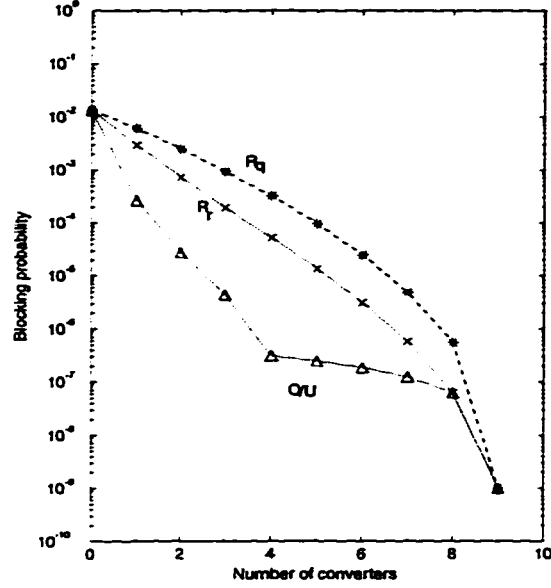


Figure 5.3: P_b vs. K for optimal/uniform (O), random (R_r) placement of K converters, and random placement of a random number of converters (R_q) with mean K , for uniform link loads ($\rho = 0.1$).

average of $q(H - 1) = K$ converters are randomly placed, and define

$$G_u^q \stackrel{\text{def}}{=} \frac{\rho_o}{\rho_q}.$$

G_u^r and G_u^q are plotted as functions of K in Figure 5.4. G_u^r can be as high as 1.6 for $K = 4$ and $P_b = 10^{-3}$. These curves were obtained by numerically solving for ρ in (5.4) and (5.5) to obtain a P_b that is within 1% of the given P_b . An approximate closed-form expression can be obtained for G_u^q for non-zero K (when $K = 0$, $G_u^q = G_u^r = 1$) by considering only the first term in (5.5) and writing

$$\begin{aligned} P_b^q(H, K) &\approx (1 - q)^{H-1} P_b^0(H) \\ &= (1 - q)^{H-1} (1 - f(H)) \\ &= (1 - q)^{H-1} (1 - \bar{\rho}^H)^F \end{aligned} \tag{5.6}$$

where $P_b^q(H, K)$ is the blocking probability of an end-to-end call when a random number of converters (with average K) are randomly placed among the $H - 1$ in-

intermediate nodes of an H -hop path, and $P_b^0(H)$ is the blocking probability of an end-to-end call on the H -hop path with no wavelength converters. This is a good approximation for $q = H/(K - 1) \leq 0.2$ (small conversion densities). In any case, $P_b^q(H, K) \geq (1 - q)^{H-1}(1 - \bar{\rho}^H)^F$ and an upper bound on ρ_q for a given P_b can be obtained by inverting (5.6). We then have

$$\begin{aligned} \rho_q &\leq 1 - \left[1 - \left(\frac{P_b}{(1 - q)^{H-1}} \right)^{\frac{1}{F}} \right]^{\frac{1}{H}} \\ &\leq \frac{-1}{H} \ln \left[1 - \left(\frac{P_b}{(1 - q)^{H-1}} \right)^{\frac{1}{F}} \right] \end{aligned} \quad (5.7)$$

where the second inequality is almost an equality for moderately large⁴ H and q small so that $\frac{P_b}{(1-q)^{H-1}}$ is not close to 1.

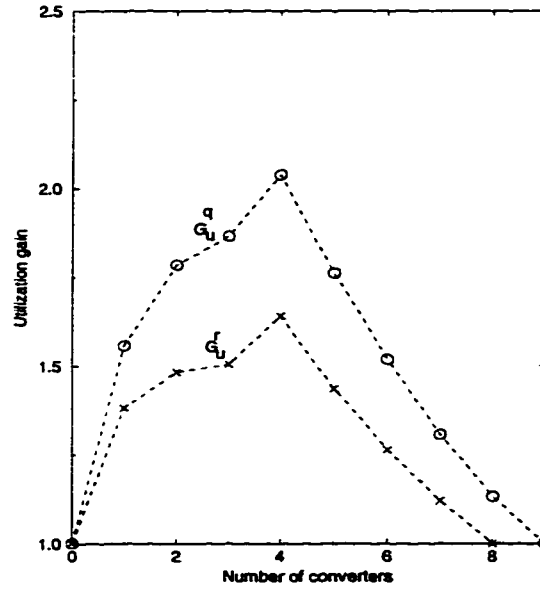


Figure 5.4: G_u^r and G_u^q vs. K for $P_b \approx 10^{-3}$.

⁴ Note that the expression is valid only if H is not too large so that $\frac{P_b}{(1-q)^{H-1}} < 1$.

Furthermore, by inverting (5.3), we have

$$\begin{aligned}
 \rho_o &= 1 - \left[1 - \left(1 - (1 - P_b)^{\frac{1}{K+1}} \right)^{\frac{1}{F}} \right]^{\frac{K+1}{H}} \\
 &\approx 1 - \left[1 - \left(\frac{P_b}{K+1} \right)^{\frac{1}{F}} \right]^{\frac{K+1}{H}} \\
 &\approx \frac{-(K+1)}{H} \ln \left[1 - \left(\frac{P_b}{K+1} \right)^{\frac{1}{F}} \right]
 \end{aligned} \tag{5.8}$$

where the first approximation in (5.8) is valid for $\frac{P_b}{K+1} \ll 1$ and the second approximation is valid for $\frac{K+1}{H} \ll 1$.

From (5.7) and (5.8), we have

$$\begin{aligned}
 G_u^q &\geq \frac{1 - \left[1 - \left(1 - (1 - P_b)^{\frac{1}{K+1}} \right)^{\frac{1}{F}} \right]^{\frac{K+1}{H}}}{1 - \left[1 - \left(\frac{P_b}{(1-q)^{H-1}} \right)^{\frac{1}{F}} \right]^{\frac{1}{H}}} \\
 &\approx \frac{(K+1) \ln \left[1 - \left(\frac{P_b}{K+1} \right)^{\frac{1}{F}} \right]}{\ln \left[1 - \left(\frac{P_b}{(1-q)^{H-1}} \right)^{\frac{1}{F}} \right]}
 \end{aligned} \tag{5.9}$$

where the approximation is valid for small P_b , moderately large H (≥ 10), and small $\frac{K+1}{H}$. G_u^q is the utilization gain when K converters are optimally placed relative to a random number of converters randomly placed, whereas G_u^r is the gain when K converters are optimally placed relative to a random placement of K converters. However, closed-form expressions for G_u^r are not possible and therefore we use the above approximation to study the behavior of the utilization gain versus F and H . As can be seen from Figure 5.4, G_u^q is an overestimate of the gain G_u^r . Notice, however, that what we have in (5.9) is a lower bound to G_u^q and this bound approximates G_u^r better for reasonably low values of q (≤ 0.2).

We plot the bound for G_u^q as a function of F for $H = 5, 10$, and 20 , and for $K = \lceil 0.2(H-1) \rceil$ in Figure 5.5. Observe that the utilization bound increases with F for all values of H . However, the rate of increase decreases as F increases. Also

notice that the bound increases with H for a given F ; this points out the importance of optimal converter placement for large F and large H . Note that the bound becomes less accurate as H increases so that a higher utilization gain than that depicted in the graph would be expected for $H = 20$ than for $H = 5$ and 10.

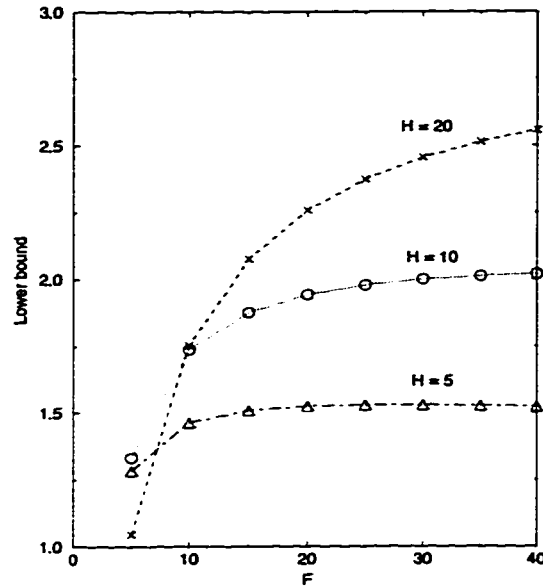


Figure 5.5: The bound for G_u^q vs. F for $P_b = 10^{-3}$ and $K = \lceil 0.2(H - 1) \rceil$.

5.2.2 Non-uniform Link Loads

We next consider the case when the link loads $\{\rho_i\}$ are not all equal. Clearly, the optimal placement is, in general, non-uniform in this case. Before we provide an optimal solution, we make a straightforward modification of (5.4) to obtain an expression for random placement.

Random Placement

We change some of the earlier definitions as follows. Define the subpath from node i to node j comprising links $i, i + 1, \dots, j - 1$ to be the *chain* $[i, j]$. Let $S_r(i, j, k)$

denote the success probability of an end-to-end call in the chain $[i, j]$ with k converters randomly distributed among the $j - 1 - i$ intermediate nodes. Then, a recurrence relation for $S_r(i, j, k)$ can be obtained as

$$S_r(i, j, k) = \begin{cases} f(i, j), & k = 0, \\ \sum_{l=i+1}^{j-k} S_r(i, l, 0) S_r(l, j, k-1) \frac{\binom{j-l-1}{k-1}}{\binom{j-i-1}{k}}, & k = 1, 2, \dots, j-1-i, \end{cases} \quad (5.10)$$

where

$$f(i, j) = 1 - \left(1 - \prod_{m=i}^{j-1} \rho_m\right)^F \quad (5.11)$$

is the success probability in the subsegment⁵ from node i to node j . Note that this subsegment contains the links $i, i+1, \dots, j-1$. The success probability of an end-to-end call when K converters are randomly placed on the path is simply given by $S_r(0, H, K)$.

Optimal Placement

When the link loads are not equal, there does not appear to be a closed-form expression for the segment lengths as in the uniform link load case. We obtain a solution based on dynamic programming along the lines of [44] for this case. In [44], the problem of placing erasure nodes in DQDB networks so that slot reuse is maximized is considered, and optimal solutions based on dynamic programming are presented. The nature of the objective function in our case is different from theirs; however, their approach can be used to solve the problem at hand, as we will see below.

For any integer m , define a *converter placement vector* $\mathbf{a} = (a_1, a_2, \dots, a_m)$ with $0 < a_i < a_{i+1} \leq H$, $i = 1, \dots, m-1$. The entries of $\mathbf{a}$ denote the placement of converters among the nodes $1, 2, \dots, H$, i.e., a_i is the location of the i th converter. The reason for the inclusion of node H as a possible converter location will be apparent below. Also define $\beta(j, \mathbf{a})$ to be the probability that an end-to-end call is successful

⁵ According to our terminology here, a chain may contain converters while a segment does not.

in chain $[0, j]$ when $\mathbf{a}$ is the placement vector, and define the set of all converter placement vectors with the m th converter at node j as

$$\Theta(m, j) \stackrel{\text{def}}{=} \{\mathbf{a} \in Z_+^m : 0 < a_i < a_{i+1} < a_m = j, i = 1, \dots, m-2\}.$$

Then $\gamma(m, j) \stackrel{\text{def}}{=} \max_{\mathbf{a} \in \Theta(m, j)} \beta(j, \mathbf{a})$ denotes the success probability in the chain $[0, j]$ with the first $m-1$ converters optimally placed among nodes $1, \dots, j-1$, and the m th converter at node j . Let us also define

$$\Gamma(m, j, k) \stackrel{\text{def}}{=} \{\mathbf{a} \in \Theta(m, j) : a_{m-1} = k\}.$$

We can obtain $\gamma(m, j)$ using the following dynamic programming procedure.

$$\begin{aligned} \gamma(m, j) &= \max_{\mathbf{a} \in \Theta(m, j)} \beta(j, \mathbf{a}) \\ &= \max_{m-1 \leq i \leq j-1} \max_{\mathbf{a} \in \Gamma(m, j, i)} \beta(j, \mathbf{a}) \\ &= \max_{m-1 \leq i \leq j-1} \max_{\mathbf{a} \in \Gamma(m, j, i)} \beta(i, \mathbf{a}) f(i, j) \\ &= \max_{m-1 \leq i \leq j-1} \left[\max_{\mathbf{a} \in \Gamma(m, j, i)} \beta(i, \mathbf{a}) \right] f(i, j) \\ &= \max_{m-1 \leq i \leq j-1} \gamma(m-1, i) f(i, j) \end{aligned} \tag{5.12}$$

for $2 \leq m \leq K+1$, $m \leq j \leq H$. Clearly $\gamma(1, j) = f(0, j)$ for $1 \leq j \leq H$.

The above recurrence relation directly leads to a $O(H^2 K)$ dynamic programming algorithm for obtaining the optimal placement vector. The optimal placement vector is the solution obtained when $K+1$ converters are placed among nodes $1, 2, \dots, H$, such that the $(K+1)$ th converter is at node H , i.e., $\mathbf{a}^{\text{opt}} = \arg \max_{\mathbf{a} \in \Theta(K+1, H)} \beta(H, \mathbf{a})$. The resulting success probability is $\gamma(K+1, H)$.

The end-to-end blocking probability is plotted against the converter position for $K=1$ for two different link load patterns in Figure 5.6. The top curve is for uniform link loads and $\rho = 0.1$, and the bottom curve is for linearly increasing link loads from link 0 to link 9. We take $\rho_0 = 0.05$, $\rho_9 = 0.1$, and $\rho_i = \rho_{i-1} + c$, $i = 1, \dots, 8$, where $c = \frac{\rho_9 - \rho_0}{9}$. Intuitively the optimal converter location should shift to the right for the

linear traffic pattern and this is indeed what is observed in the figure. However, it cannot be determined *a priori* that the converter is optimally placed at node 6 for this traffic pattern. From the figure, note that converter placement in a non-optimal position would increase P_b by a factor of 2 at least and the performance degradation could be as high as two orders of magnitude.

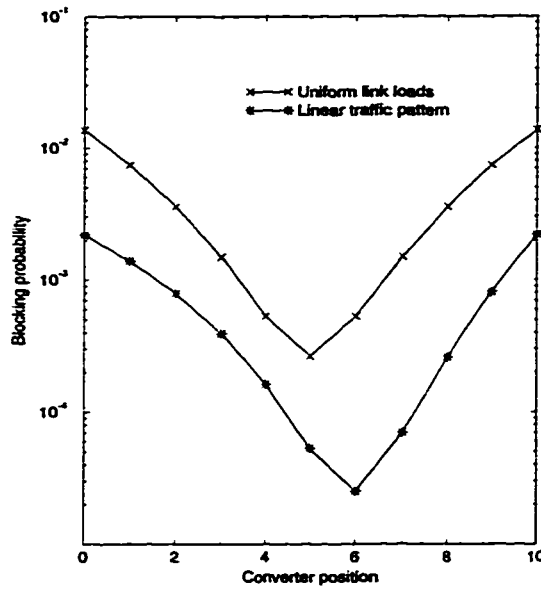


Figure 5.6: P_b vs. the converter position for $K = 1$ for uniform link loads ($\rho = 0.1$) and for linearly increasing link loads ($\rho_0 = 0.05, \rho_9 = 0.1$).

The performance of random (R), uniform (U), and optimal (O) placements are compared for a linear link load traffic pattern in Figure 5.7. The uniform placement performs quite poorly for a large number of converters. This is due to the fact that converters are best placed towards the end of the path whereas the uniform placement policy would place some towards the beginning of the path, where they are expected to be less useful. The optimal placement results in dramatically better performance than random placement.

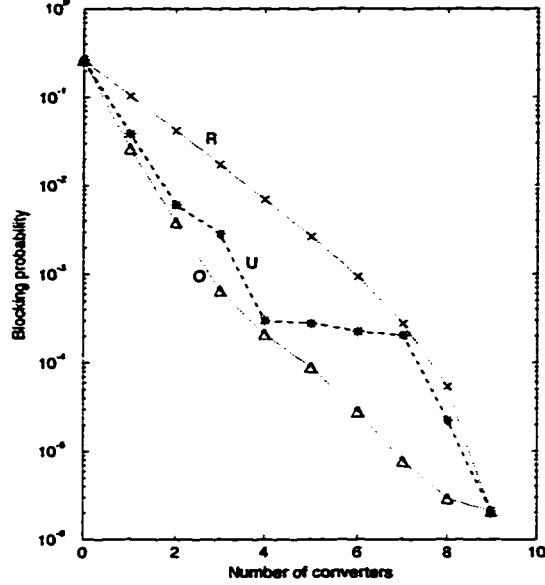


Figure 5.7: P_b vs. K for random (R), uniform (U) and optimal (O) placements. Link loads are linearly increasing with $\rho_0 = 0.125, \rho_9 = 0.25$.

5.2.3 Placement for Optimal Average Performance

Until this point, we have considered the placement problem from the point of view of a connection request from node 0 to node H . However, when a path is part of a large network, it may be of interest to minimize the average blocking probability over all traffic using the path, rather than just the end-to-end traffic. We now modify the solution above to take this into account. Random placement is considered first.

Random Placement

Let $\alpha(m, j, s, d)$ denote the probability that a call from s to d will be successful in the chain $[0, j]$ when m converters are placed in $[1, j]$ with the m th converter at node j and the other $m - 1$ converters randomly distributed in $[1, j - 1]$. Define $\delta(j, \mathbf{a}, s, d)$ to be the probability that an (s, d) call succeeds in the chain $[0, j]$ when $\mathbf{a}$ is the

placement vector. Then,

$$\alpha(m, j, s, d) = E_{\mathbf{a} \in \Theta(m, j)} \delta(j, \mathbf{a}, s, d),$$

where E_x denotes the expectation over x . Note that all $\mathbf{a} \in \Theta(m, j)$ are equiprobable with random converter placement. Then, similar to (5.10), we can write

$$\alpha(m, j, s, d) = \begin{cases} h(0, j, s, d), & m = 1, \\ \sum_{i=m-1}^{j-1} \alpha(m-1, i, s, d) h(i, j, s, d) \frac{\binom{i-1}{m-2}}{\binom{j-1}{m-1}}, & m = 2, \dots, j, \end{cases} \quad (5.13)$$

where $h(i, j, s, d)$ is the probability that an (s, d) call is successful in segment $[i, j]$. $h(i, j, s, d)$ is given by

$$h(i, j, s, d) = \begin{cases} f(s, d), & i \leq s < d \leq j, \\ f(i, j), & s \leq i, d \geq j, \\ f(i, d), & s \leq i, i+1 \leq d \leq j-1, \\ f(s, j), & i \leq s \leq j-1, d \geq j, \\ 1, & d \leq i, s \geq j, \end{cases} \quad (5.14)$$

where f is as defined in (5.11). The average success probability considering all (s, d) calls is then $E_{(s,d)} \alpha(K+1, H, s, d)$.

Optimal Placement

The optimal placement vector $\mathbf{a}^{\text{opt}}$ that minimizes the average call blocking probability is given by

$$\begin{aligned} \mathbf{a}^{\text{opt}} &= \arg \min_{\mathbf{a} \in \Theta(K+1, H)} E_{(s,d)} [1 - \delta(H, \mathbf{a}, s, d)] \\ &= \arg \max_{\mathbf{a} \in \Theta(K+1, H)} -E_{(s,d)} [1 - \delta(H, \mathbf{a}, s, d)]. \end{aligned}$$

Note that the max and the E operators cannot be interchanged, and as a result, it is not possible to obtain a recurrence relation such as the one in (5.12). However, an

approximate solution to this problem can be obtained as follows. Instead of maximizing $E_{(s,d)}\delta(H, \mathbf{a}, s, d)$, we will obtain the placement configuration that maximizes $E_{(s,d)} \ln \delta(H, \mathbf{a}, s, d)$. Note that $1 - \delta(H, \mathbf{a}, s, d)$ is the probability that the (s, d) call is *blocked* in the chain $[0, H]$ when the placement vector is $\mathbf{a}$. For reasonably low blocking probabilities ($< 10^{-2}$), because of the fact that $\ln x \approx x - 1$ when $x \approx 1$, $E_{(s,d)} \ln \delta(H, \mathbf{a}, s, d)$ is an excellent approximation for $-E_{(s,d)}(1 - \delta(H, \mathbf{a}, s, d))$ which is the original objective function to maximize.

Define $\zeta(j, \mathbf{a})$ to be

$$\zeta(j, \mathbf{a}) \stackrel{\text{def}}{=} E_{(s,d)} \ln \delta(j, \mathbf{a}, s, d) \quad (5.15)$$

Now if $\xi(m, j) \stackrel{\text{def}}{=} \max_{\mathbf{a} \in \Theta(m, j)} \zeta(j, \mathbf{a})$, then

$$\begin{aligned} \xi(m, j) &= \max_{m-1 \leq i \leq j-1} \max_{\mathbf{a} \in \Gamma(m, j, i)} \zeta(j, \mathbf{a}) \\ &= \max_{m-1 \leq i \leq j-1} \left\{ \max_{\mathbf{a} \in \Gamma(m, j, i)} \left[E_{(s,d)} \ln \delta(i, \mathbf{a}, s, d) + E_{(s,d)} \ln h(i, j, s, d) \right] \right\} \\ &= \max_{m-1 \leq i \leq j-1} \left\{ \max_{\mathbf{a} \in \Gamma(m, j, i)} \left[\zeta(i, \mathbf{a}) + E_{(s,d)} \ln h(i, j, s, d) \right] \right\} \\ &= \max_{m-1 \leq i \leq j-1} \left\{ \left[\max_{\mathbf{a} \in \Gamma(m, j, i)} \zeta(i, \mathbf{a}) \right] + E_{(s,d)} \ln h(i, j, s, d) \right\} \\ &= \max_{m-1 \leq i \leq j-1} \left\{ \xi(m-1, i) + E_{(s,d)} \ln h(i, j, s, d) \right\} \end{aligned} \quad (5.16)$$

for $2 \leq m \leq K+1$, $m \leq j \leq H$, and $\xi(1, j) = E_{(s,d)} \ln h(0, j, s, d)$ for $1 \leq j \leq H$.

The optimal placement $\mathbf{a}^{\text{opt}}$ is now given by $\arg \max_{\mathbf{a} \in \Theta(K+1, H)} \xi(K+1, H)$ and (5.16) defines a recurrence relation that can be solved via dynamic programming as before. In Figure 5.8, we plot the average blocking probability for optimal and random placements and the following traffic pattern. The load per wavelength between each (s, d) pair, $\lambda_{sd} = 0.01$, $d > s$, so that the load per wavelength on link i , $\rho_i = 0.01(i+1)(H-i)$. In this traffic pattern, the link loads increase from link 0 linearly until link 4 and then decrease linearly from link 5 to link 9. We also plot the performance for an end-to-end call in the figure for comparison. We observe that optimal placement provides considerable improvement in average performance as well, relative to random

placement. For instance, to achieve an average $P_b \approx 10^{-3}$, 3 converters suffice if placed optimally while 7 converters would be required if placed randomly.

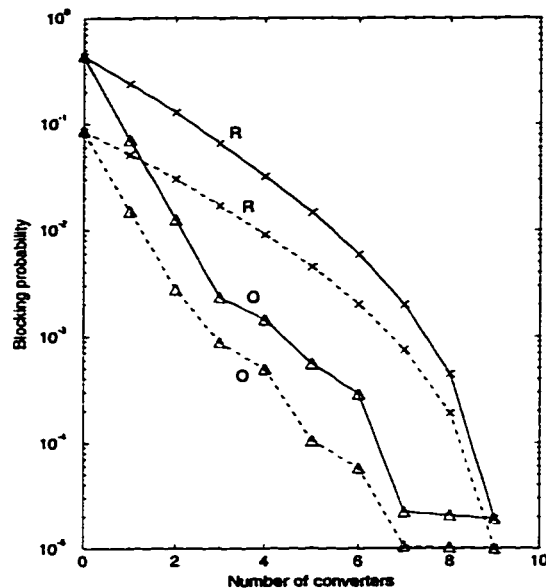


Figure 5.8: P_b vs. K for random (R), and optimal (O) placements. $H = 10$, $F = 10$. $\rho_i = 0.01(H - i)(i + 1)$. Solid lines: End-to-end performance, Dashed lines: Average performance.

5.3 Converter Placement in a Bus

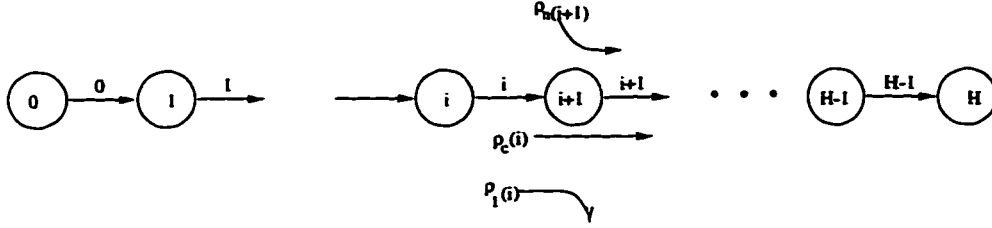
In the previous section, we considered the problem of optimally placing a given number of converters on a path assuming that link loads are independent. In this section, we show how that framework can be applied to a unidirectional bus network where there is no interfering traffic [18] and the traffic in the network is solely due to the traffic originating at the nodes of the bus. In such a scenario, the link load independence assumption of the previous section is not appropriate. We show in this section how our dynamic programming solutions are applicable even in the presence of link load correlation.

A key step in the derivation of the recurrence relation in (5.16) is the third equality where we have assumed that the success probability of an (s, d) call in chain $[0, i]$ is independent of the success probability of the call in segment $[i, j]$ when node i is a converter. This assumption is still needed; that is, we continue to assume that the success probabilities of a call in disjoint segments are statistically independent. However, note that we do not need the independence assumption between the links of the same segment. Since the fraction of the nodes with converters is typically small, the most important effects of link load correlation will be taken into account in our formulation. Furthermore, the effects of load correlation between two consecutive links on a path separated by a converter are not as significant as when there is no converter between the links. This is because, when there is no converter, it is not only the dependence in the number of wavelengths on the two links that is significant, but also the actual wavelengths used on those two links. Note that this latter dependence is irrelevant if there is a converter between the two links.

The performance model we use here is the one proposed by Barry in [18, 20]. For completeness, we present the relevant details of Barry's model here. The only quantity of interest to us here is the probability of success $f(i, j)$ in a subsegment $[i, j]$ since all other quantities are defined in terms of it.

Consider the bus in Figure 5.9 and let $\rho_n(i)$ be the load per wavelength that enters the network at node i , $\rho_l(i)$ the load per wavelength that leaves the network at node $i + 1$, and $\rho_c(i)$ the load per wavelength that continues through node $i + 1$ and uses links i and $i + 1$. Also, let us denote the load per wavelength on link i by $\rho(i) (= \rho_n(i) + \rho_c(i - 1))$.

Since wavelengths are considered to be independent of each other in this model, it suffices to look at a single wavelength, say w_1 . Then, $P_l(i) \stackrel{\text{def}}{=} \rho_l(i)/\rho(i)$ is defined as the probability that a call on wavelength w_1 uses link i and leaves at node $i + 1$. Also, let $P_n(i)$ be the probability that a new call enters the network at node i and uses link i on wavelength w_1 given that w_1 is not used by another call on link i . $\rho(i)$

Figure 5.9: A bus of length H .

is then given by [18]

$$\rho(i) = \rho(i-1) [1 - P_l(i-1) + P_l(i-1)P_n(i)] + (1 - \rho(i-1))P_n(i)$$

and, therefore,

$$P_n(i) = \frac{\rho(i) - \rho(i-1)[1 - P_l(i-1)]}{1 - \rho(i-1)[1 - P_l(i-1)]}.$$

The probability that at least one wavelength is available on all links of a subsegment $[i, j]$ is now found as

$$\begin{aligned} f(i, j) &= 1 - [1 - P(w_1 \text{ free on link } i) \\ &\quad \cdot \prod_{k=i+1}^{j-1} P(w_1 \text{ free on link } k \mid w_1 \text{ free on link } k-1)]^F \\ &= 1 - [1 - (1 - \rho(i)) \prod_{k=i+1}^{j-1} (1 - P_n(k))]^F. \end{aligned} \quad (5.17)$$

Given the traffic matrix $\Lambda = [\lambda_{sd}]$ where λ_{sd} is the load per wavelength between nodes s and d , $\rho_l(i)$, $\rho_n(i)$, $\rho_c(i)$, and $\rho(i)$ are calculated in the following obvious way.

$$\begin{aligned} \rho(i) &= \sum_{s=0}^i \sum_{d=i+1}^H \lambda_{sd} \\ \rho_l(i) &= \sum_{s=0}^i \lambda_{s,i+1} \\ \rho_n(i) &= \sum_{d=i+1}^H \lambda_{id} \\ \rho_c(i) &= \rho(i) - \rho_l(i). \end{aligned}$$

As before, (5.13) and (5.16) are used to evaluate the performance under random and optimal placement, respectively. However, (5.17) is used instead of (5.11) for computing the success probability in a subsegment, thus taking the link load correlation in the bus topology into account.

We show P_b as a function of K for a 10-node bus with $F = 10$ in Figure 5.10. The link loads are kept constant at $\rho = 0.4$. One way of achieving this is by varying the station loads $\rho_n(i)$ and assuming that a call arriving at a node can have any downstream node as its destination with equal probability. Then, $\rho_n(0) = \rho$, and $\rho_n(i) = \rho - \sum_{j=0}^{i-1} \frac{H-i}{H-j} \rho_n(i-1)$, $i = 1, 2, \dots, H-1$, and $\lambda_{ij} = \frac{\rho_n(i)}{H-i}$ for $j > i$, and 0 otherwise. Here, we observe that the performance improvement over random placement due to optimal converter placement is not as high as in a path with independent link loads. However, the improvement is significant considering the fact that the performance does not improve considerably as the *number* of converters increases. We have only considered one traffic pattern here and there may be other traffic patterns producing uniform link loads which could cause a more significant performance difference between optimal and random placements.

Notice further that uniform converter placement is no longer optimal even though link loads are uniform, because of the link load correlation as well as the fact that we have considered the optimal placement with respect to all traffic and not just end-to-end traffic. P_b vs. the converter position is plotted in Figure 5.11 for $K = 1$, and a uniform link load of $\rho = 0.3$. It is certainly not clear beforehand that the optimal solution would place the converter at node 7. This is in fact surprising considering the fact that node 7 is neither the node with the highest transit (continuing) traffic, nor is the node mixing the highest amount of traffic (high ρ_n and ρ_l). ρ_n and ρ_c for each node are plotted in Figure 5.12 for this traffic pattern⁶. For $K = 2$, the optimal placements turn out to be at nodes 5 and 8, again seemingly unintuitive.

⁶ By our previous definition, ρ_c is the traffic that continues through node $i + 1$; however, in the figure ρ_c refers to the traffic continuing through the corresponding node number.

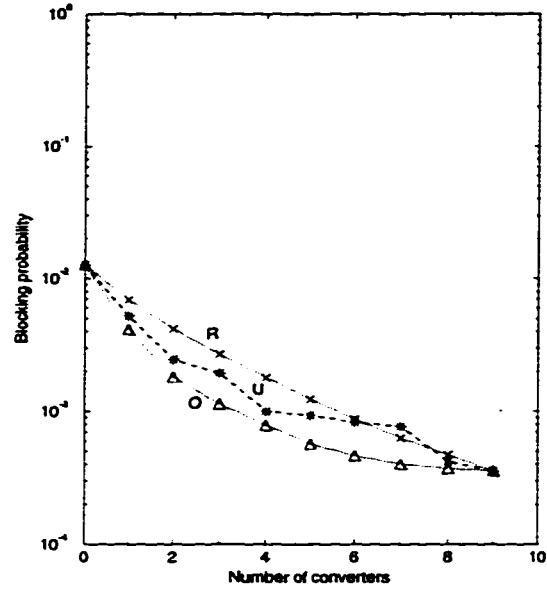


Figure 5.10: P_b vs. K for uniform link loads, $\rho = 0.4$. $H = 10$, $F = 10$.

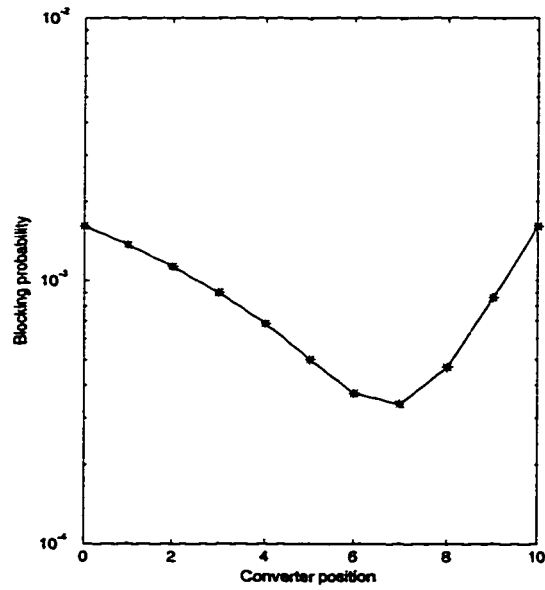


Figure 5.11: P_b vs. converter position for $K = 1$ and uniform link loads, $\rho = 0.3$. $H = 10$, $F = 10$.

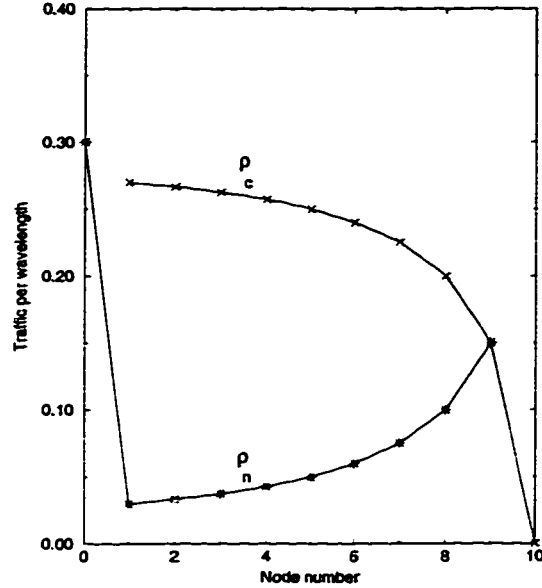


Figure 5.12: Traffic vs. node number for uniform link loads, $\rho = 0.3$.

We remark here that this formulation can be applied to a double-bus network as well. The only necessary change is that the objective function must now maximize the average success probability over both buses together. The traffic model will have to be applied to the two buses separately.

5.4 Converter Placement in a Ring

The ring is a popular optical network topology because of its simplicity [16, 22, 45]. In Chapter 3, we presented a performance model for determining the benefits of wavelength conversion in ring networks with uniform Poisson traffic. In this section, we present a dynamic programming solution for optimal converter placement in a ring topology. As in the bus, the optimality criterion is the average blocking probability. We will mainly be interested in the results obtained using the binomial model here, i.e., we continue to assume that all wavelengths are used on a link with the same probability and wavelength events are independent of each other. However, the

correlation between wavelength usage events on consecutive links will be taken into account as in the previous section.

Consider a unidirectional ring network with H nodes, numbered from 0 to $H - 1$, where the link from node i to node $(i + 1) \bmod H$ is labeled as link i . The optimal placement for the ring is obtained by using the fact that if a converter is placed at a node, say i , then because of the assumed independence of the success probability of a call in disjoint segments, the ring can be logically broken at node i to create two nodes i' and i'' and creating a bus of length H (see Figure 5.13). Let the left end-node of the resulting bus be i' and the right end-node i'' . The traffic matrix Λ for the ring is modified to form a traffic matrix Λ' for the bus as follows.

$\lambda'_{si''} = \lambda'_{i'd} = \lambda_{sd}$ if chain $[s, d]$ contains node i as an intermediate node in the ring and $\lambda'_{sd} = \lambda_{sd}$ otherwise.

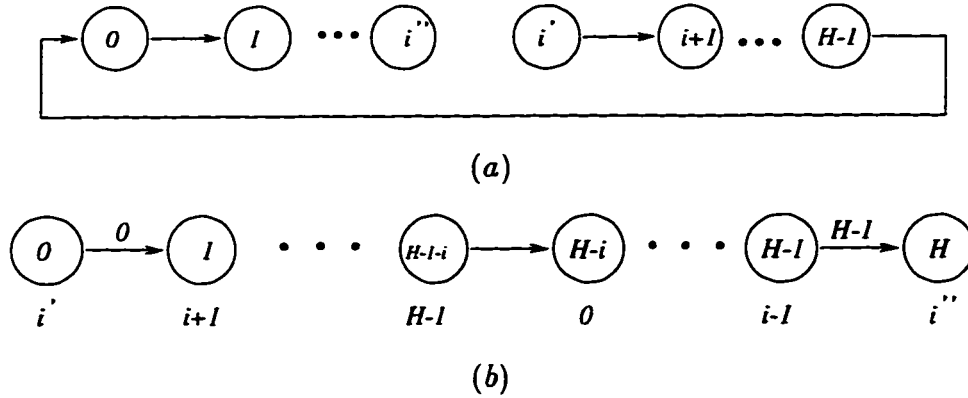


Figure 5.13: (a) The ring logically broken at node i where a converter is assumed to be placed. (b) The resulting bus after relabeling (corresponding old labels are shown below the nodes).

Given an H -node ring and K converters, the average success probabilities can be obtained by conditioning on the lowest indexed node i having a converter, $i = 0, 1, \dots, H - K$. The ring is then broken at node i as described above and the optimal placement of $K - 1$ converters on the logical bus given i is the lowest indexed

node with a converter is obtained as in the previous section. The optimal placement over the ring is the placement that maximizes the success probability over all i , and the performance with random placement is the performance averaged over all i . This is stated formally below.

Suppose for ease of notation that the nodes and links of the bus derived from the ring are relabeled so that node i' is node 0, node i'' is node H , and the link out of the relabeled node i is link i . Note that there are no converters that are placed in nodes $H - i$ through $H - 1$ of the bus since node i was assumed to be the lowest numbered node in the ring with a converter. The ring broken at node i , and the corresponding bus after relabeling are shown in Figure 5.13(a) and 5.13(b), respectively. The original labels of the ring are shown below the nodes in Figure 5.13(b).

5.4.1 Random Placement

To obtain the performance with random placement, define

$$\kappa(m, j) \stackrel{\text{def}}{=} E_{\mathbf{a} \in \Theta(m, j)} \zeta(j, \mathbf{a})$$

where $\zeta(j, \mathbf{a})$ is as defined in (5.15), to be the average over all (s, d) calls of the logarithm of the success probability in $[0, j]$ when $m - 1$ converters are randomly placed in nodes $1, 2, \dots, j - 1$, and the m th converter is placed at node j of the relabeled bus. Let

$$\Upsilon(K, H, t) \stackrel{\text{def}}{=} \{\mathbf{a} \in \Theta(K, H) : a_{K-1} \leq t\}$$

denote the set of placement vectors which place converter K at node H and converter $K - 1$ at a node whose index is no more than t .

Then, under random converter placement, we have

$$\begin{aligned} -P_b &\approx E_i E_{\mathbf{a} \in \Upsilon(K, H, H-1-i)} \zeta(H, \mathbf{a}) \\ &= E_i E_{\mathbf{a} \in \Upsilon(K, H, H-1-i)} E_{(s, d)} \ln \delta(H, \mathbf{a}, s, d) \end{aligned}$$

$$\begin{aligned}
&\stackrel{\text{a}}{=} E_i \left\{ \sum_{k=K-1}^{H-1-i} \frac{\binom{k-1}{K-2}}{\binom{H-1-i}{K-1}} E_{\mathbf{a} \in \Theta(K-1,k)} E_{(s,d)} \{ \ln \delta(k, \mathbf{a}, s, d) + \ln h(k, H, s, d) \} \right\} \\
&= E_i \left\{ \sum_{k=K-1}^{H-1-i} \frac{\binom{k-1}{K-2}}{\binom{H-1-i}{K-1}} \left[\kappa(K-1, k) + E_{(s,d)} \ln h(k, H, s, d) \right] \right\} \\
&\stackrel{\text{b}}{=} \frac{\binom{H-1-i}{K-1}}{\binom{H}{K}} \left\{ \sum_{k=K-1}^{H-1-i} \frac{\binom{k-1}{K-2}}{\binom{H-1-i}{K-1}} \left[\kappa(K-1, k) + E_{(s,d)} \ln h(k, H, s, d) \right] \right\} \\
&= \frac{1}{\binom{H}{K}} \left\{ \sum_{k=K-1}^{H-1-i} \binom{k-1}{K-2} \left[\kappa(K-1, k) + E_{(s,d)} \ln h(k, H, s, d) \right] \right\} \quad (5.18)
\end{aligned}$$

where $\kappa(m, j)$ is obtained, similar to α in (5.13), to be

$$\kappa(m, j) = \begin{cases} E_{(s,d)} \ln h(0, j, s, d), & m = 1, \\ \sum_{l=m-1}^{j-1} \frac{\binom{l-1}{m-2}}{\binom{j-1}{m-1}} \left\{ \kappa(m-1, l) + E_{(s,d)} \ln h(l, j, s, d) \right\}, & m = 2, \dots, j. \end{cases}$$

In Step (a) of (5.18), $\binom{k-1}{K-2} / \binom{H-1-i}{K-1}$ is the probability that the $(K-1)$ th converter is placed at node k given that $K-1$ converters are placed randomly among nodes $1, 2, \dots, H-1-i$. In Step (b), $\binom{H-1-i}{K-1} / \binom{H}{K}$ is the probability that the node with the lowest index in the ring that has a converter is node i , given that K converters are placed randomly. We also have used the definitions of ζ , κ , and h in (5.18). Note that we did not employ the logarithm approximation in computing the performance for random placement in Section 5.2. In the ring, however, because the traffic through a node is split into two streams, the logarithm approximation is necessary in Step (a) of (5.18).

5.4.2 Optimal Placement

Along similar lines, the optimal placement vector is obtained as

$$\mathbf{a}^{\text{opt}} = \arg \max_i \max_{\mathbf{a} \in \Upsilon(K, H, H-1-i)} \zeta(H, \mathbf{a}).$$

Now,

$$\max_i \max_{\mathbf{a} \in \Upsilon(K, H, H-1-i)} \zeta(H, \mathbf{a})$$

$$\begin{aligned}
&= \max_i \max_{K-1 \leq k \leq H-i-1} \max_{a \in \Theta(K-1, k)} \{E_{(s,d)} \ln \delta(k, a, s, d) + \ln h(k, H, s, d)\} \\
&= \max_i \max_{K-1 \leq k \leq H-i-1} \{\xi(K-1, k) + E_{(s,d)} \ln h(k, H, s, d)\}
\end{aligned}$$

and $\xi(m, j)$ is given by (5.16). This provides the required solution via dynamic programming as before.

5.4.3 Results

We consider a 10-node ring network with $F = 10$. Figure 5.14 shows P_b plotted against K for optimal/uniform and random placement of converters when the link loads are uniform (as are loads between any (s, d) pair, $s \neq d$) with $\rho = 0.4$. Again, due to the considerable load correlation, the performance improvement of optimal placement over random placement is only marginal compared to the gains achievable in a path where there is insignificant load correlation.

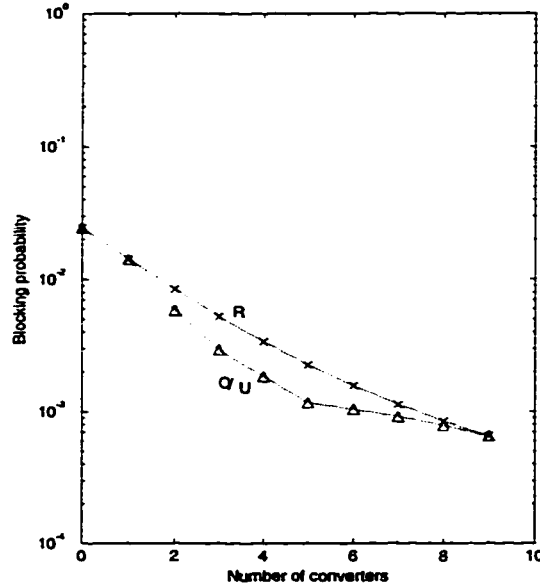


Figure 5.14: P_b vs. K for uniform link loads, $\rho = 0.4$. $H = 10$, $F = 10$.

However, optimal placement could assume importance when the traffic is not uniform. For instance, if all traffic were local traffic (from a node to its neighbor) but

for some traffic that goes from node 0 to node $H - 1$, then the situation is similar to the one in the path with no load correlation. In that case, we observed that optimal placement can provide a significant boost in end-to-end blocking performance (see Figure 5.3). Note that converters do not affect the blocking performance of local calls.

5.5 *The Effect of Traffic Model on Converter Placement*

In this section, we discuss the effects of the traffic model used in obtaining the blocking probabilities for random and optimal placements. As mentioned earlier, our dynamic programming solution techniques do not assume a specific traffic model even though the results we have presented thus far are for the binomial model.

The effect of the traffic model appears in our solutions in the form of the probability $f(i, j)$, the success probability of a call in segment $[i, j]$. In the binomial model, this probability was given by (5.11) and (5.17) when link loads are independent and strongly correlated, respectively. This probability can also be obtained via the models presented in [22] and Chapter 3 for Poisson traffic, when link loads are independent and correlated, respectively. Note that those models were presented for uniform link loads; when the link loads are non-uniform, the two-hop model presented therein will need to be applied to every pair of consecutive links in the network. We omit the details here and present only a comparison of a sampling of the results obtained using those models and the binomial model we have used so far in this chapter. We proved in Section 5.2 that in a path with uniform link loads and the link load independence assumption, uniform placement is the optimal placement for the end-to-end traffic. This was achieved by showing that the function $f(l)$ in (5.1) was a concave function of l . When the traffic model is Poisson, it appears difficult to prove the concavity still holds, but the function appears to be concave for a wide range of the parameters ρ , F , and l .

In Figure 5.15, we plot end-to-end P_b against K for a path with $H = 10$, $F = 10$, and the load per wavelength per fiber, $\rho = 0.1$, with link load independence and the Poisson traffic model. Comparing this with Figure 5.3, we see that the numerical values of P_b are higher for the Poisson model⁷, but the conclusions regarding converter placement are the same. For example, the performance difference between optimal and random placement increases with K until it reaches a maximum at $K = 4$, and then starts decreasing, in both cases.

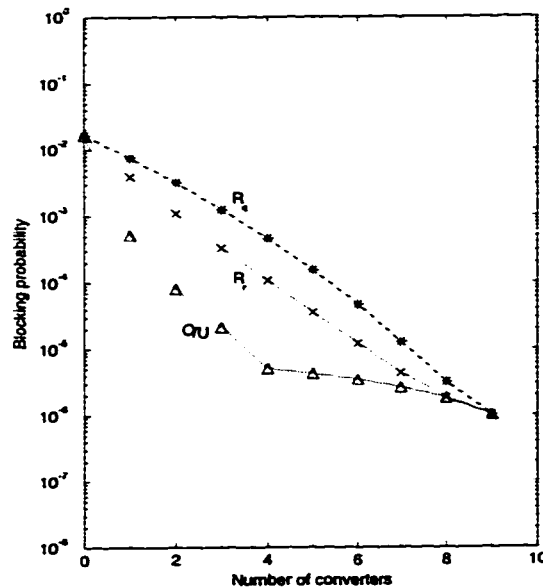


Figure 5.15: P_b of an end-to-end call vs. K for a 10-hop path with optimal/uniform (O), random (R_r) placement of K converters, and random placement of a random number of converters (R_q) with mean K . Poisson traffic model.

This is also evident from Figure 5.16 where P_b is plotted against the converter position for uniform link loads and a linear link load traffic pattern when $K = 1$. Contrast this with Figure 5.6 and observe that the relative performance over the

⁷ This is due to the fact that the busy wavelength distribution on a link offered Poisson traffic has a heavier tail than the binomial distribution for the same load (when loads are low enough so that effects of blocking on carried load can be ignored).

position of the converter is the same irrespective of the traffic model. This result, as well as related numerical experimentation, suggests that the optimal placement is largely insensitive to the actual traffic model.

However, the binomial and the Poisson model do not always produce the same optimal placement. We have been able to find some cases where the two models produce different placements, but these tend to happen when the blocking probabilities are high (> 0.1).

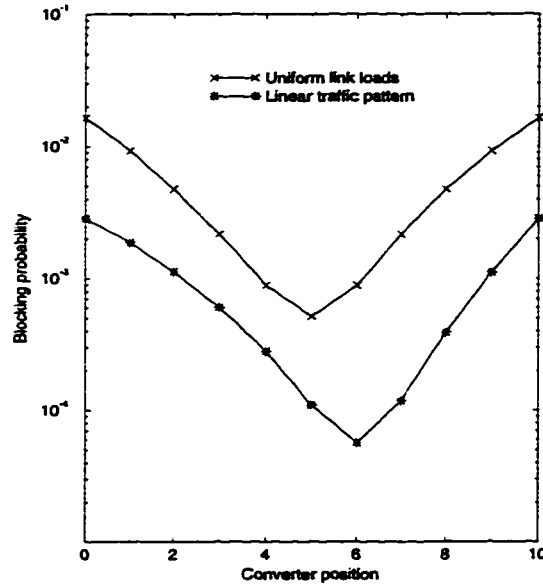


Figure 5.16: P_b vs. the converter position for $K = 1$ for uniform link loads ($\rho = 0.1$) and for linearly increasing link loads ($\rho_0 = 0.05, \rho_9 = 0.1$). Poisson traffic model.

5.6 Conclusions

In this chapter, we considered the problem of optimally placing a given number of converters in a path that is part of a dense network, in bus topologies, and in ring topologies. We first showed that uniformly spaced converters produce optimal performance for an end-to-end call on a path when the link loads are uncorrelated

and uniform. When the link loads are non-uniform or when other calls are considered, we provided solutions based on dynamic programming for the optimal placement of converters. Expressions for blocking probability when the converters are randomly placed were also derived.

The results indicate that optimal converter placement is an important problem, especially when there is negligible correlation between link loads. Our solutions are applicable to a variety of traffic models, even though most of the results presented here were for the binomial model.

In Section 5.5, we mentioned that the optimal placement is largely insensitive to the traffic model, but does depend on it at high blocking probabilities. We remark here that the optimal placement is not only somewhat dependent on the traffic model but also on the actual load when a given traffic pattern is scaled. That is, the optimal placements may change as the load is scaled up with the same traffic pattern. However, the changes do not appear to be arbitrary and there seem to be well-defined regions where the optimal placement configurations remain the same. This phenomenon needs further investigation.

Finally, the problem of optimally placing converters in an arbitrary topology is still an open problem and effective heuristics, possibly based on our solutions presented here, will be required.

Chapter 6

DYNAMIC WAVELENGTH ASSIGNMENT IN FIXED-ROUTING NETWORKS

6.1 *Introduction*

In the previous chapters, we presented performance models for wavelength-routed networks with sparse wavelength conversion, and optimal solutions for the converter placement problem. In this chapter, we consider networks with no wavelength conversion. Eliminating wavelength conversion significantly reduces the cost of the wavelength-routing switch; however it may lead to reduced network efficiency because the same wavelength must be available on each link of a route for a call to be established. For example, there may be a route which has a wavelength available on each of its links but they may not be the same wavelength, thereby causing the call to be rejected.

A good routing and wavelength assignment (RWA) algorithm is critically important to increasing the efficiency of such networks. The RWA algorithm is responsible for selecting a suitable route and wavelength among the many possible choices for establishing the call. After the algorithm has made its choice, the transmitting lasers and receivers at various nodes are tuned to the appropriate wavelengths, and the switches are set to switch the selected wavelength along the selected route. Communication can then begin. Note that if unlimited wavelength conversion is available, wavelength assignment is a trivial problem although a good routing algorithm is still needed.

The RWA problem has been extensively studied in several ways. Models for

obtaining the network blocking performance with and without wavelength conversion under dynamic Poisson traffic have been proposed in [18, 19, 20, 21, 22, 27]. The models predict the blocking probability when the random wavelength assignment algorithm is used¹. Lower bounds on the blocking probability in networks with and without conversion are derived in [10]. Several RWA algorithms that differ in the traffic assumptions and the performance metric used have also been proposed. The traffic assumptions generally fall into one of two categories: (a) static traffic, wherein a set of call requests is given and routes and wavelengths have to be assigned to calls such that some metric is optimized [24, 46, 47], and (b) dynamic traffic, wherein calls arrive to and depart from the network one by one in a random manner [21, 22, 32, 45, 48, 49]. The performance metric in this latter case is typically the call blocking probability.

In this chapter too, we consider a wavelength-routing network operating in circuit-switched mode under dynamic traffic. To study the importance of wavelength assignment in the RWA algorithm, we focus on it and assume that the routes for the calls are fixed. We propose a new wavelength assignment algorithm called the Max.Sum ($M\Sigma$) algorithm. In Section 6.2, we define the problem and present the algorithm. A very simple implementation of the algorithm for the unidirectional ring topology and a more general implementation for arbitrary topologies is given in Section 6.3. In Section 6.4, we extend the analytical model of [22, 27] to multi-fiber networks. Simulation results of the performance of our algorithm are compared with those of previously proposed algorithms on two topologies with different network connectivities – a ring and a mesh-torus, in Section 6.5. The results show that our algorithm often performs better (and no worse) for the example networks we consider. The results also show that the analytical model is accurate enough to make qualitative predictions about network performance under random wavelength assignment. Finally, we present our

¹ Analytical results for “smart” wavelength assignment algorithms are hard to obtain.

conclusions in Section 6.6.

The work to be presented in this chapter was done jointly with Dr. Richard Barry of MIT Lincoln Laboratory. He conceived the original idea of the Max_Sum algorithm for single-fiber rings, and developed the simulation programs for performance study of the ring.

6.2 *The Problem and the $M\Sigma$ Algorithm*

Assume that the network is in a certain state, i.e., a set of calls is already established and wavelengths are assigned to those calls. Suppose a new call request arrives. The problem is to assign a wavelength to the call without rearranging the already established calls so that the call blocking probability is minimized. It is extremely hard to obtain an exact solution to this problem even if the traffic statistics are accurately known [24, 32, 50]. Heuristic solutions are therefore necessary. We restrict ourselves in this chapter to the study of algorithms that admit a call if a wavelength is available, i.e., algorithms that do not perform admission control.

All existing algorithms assign wavelengths based on the *current state* of the network. Our algorithm differs from those algorithms in that we also consider the state of the network *after* the new call has been established. Our algorithm tries to choose the wavelength that leaves the network in a “good” state for future calls. The goodness of a state is measured by a new concept called the *value* of the network. The value of the network is an important single-valued metric that the network operator can use for planning. This is an added advantage of our algorithm over other proposed algorithms which do not provide a good metric of the network state for the operator.

Our algorithm is based on using path capacities to calculate the network value. In particular, each path of the network has a path capacity which is defined to be the number of calls that can be established on that path. Note that the path capacities change over time and depend upon the current network state. The network state

changes when a call is established and the next state depends on the particular wavelength that was assigned to the call. The value of the network is a function of the path capacities. Our algorithm assigns wavelengths to calls such that the next state value is maximized. We formalize this below.

6.2.1 The Algorithm

Suppose the network is in an arbitrary state ψ and there is a new desired call along path p^* . The goal is to make a wavelength choice for the call so that the network is in a good state after the call is established. This can be posed as the problem of maximizing a function $V(\alpha)$ of the resulting state α after call establishment. The maximization is done over all possible resulting states. In our algorithm, we restrict V to be functions of path capacities (formally defined below), i.e., $V(\alpha) = g([C(\alpha, p) : p \in P])$ where P is the set of all possible paths, and $C(\alpha, p)$ is the path capacity of p in state α .

Path capacities are defined in terms of links capacities. The *link capacity* of link l on wavelength f , $c(\psi, l, f)$ is defined as the number of fibers on which f is available (unused) on link l . The capacity of path p on wavelength f , $C'(\psi, p, f)$ is then defined as the number of fibers on which f is available on the most congested link along the path, i.e.,

$$C'(\psi, p, f) \stackrel{\text{def}}{=} \min_{l \in L(p)} c(\psi, l, f)$$

where $L(p)$ is the set of links comprising path p . The *path capacity* is

$$C(\psi, p) \stackrel{\text{def}}{=} \sum_{f=1}^F C'(\psi, p, f)$$

where F is the maximum number of wavelengths available on a fiber in an empty network.

Let $\Omega(\psi, p^*)$ be the set of all possible wavelengths that are available to establish the call, and let $\psi'(f)$ be the next state of the network if wavelength f is assigned to

the call. Our algorithm chooses the wavelength f^* such that:

$$f^* = \arg \max_{f \in \Omega(\psi, p^*)} V(\psi'(f)).$$

Each function V gives a different algorithm. Some possible functions that can be used are:

1. $V(\alpha) = \sum_{p \in P} C(\alpha, p)$, called the Max_Sum ($M\Sigma$) algorithm,
2. $V(\alpha) = \sum_{p \in P} w(p)C(\alpha, p)$, called the Max_Weighted_Sum ($MW\Sigma$) algorithm,
and
3. $V(\psi, \alpha) = \min_{p \in P} C(\alpha, p)$, called the Max_Min algorithm.

In this chapter, we only present the $M\Sigma$ algorithm and its performance results. In the next section, we discuss the $M\Sigma$ algorithm in more detail and present a simple implementation of the algorithm for a widely used optical network topology, the ring, and a more general implementation that is applicable to arbitrary mesh topologies. As we will see, the algorithm outperforms previously proposed algorithms in many cases.

6.3 Implementing the $M\Sigma$ Algorithm

As mentioned earlier, the $M\Sigma$ algorithm chooses f to maximize $\sum_{p \in P} C(\psi'(f), p)$. It turns out to be easier to compute a cost function ($= \sum_{p \in P} (C(\psi, p) - C(\psi'(f), p))$) and minimize the cost than to compute $\sum_{p \in P} C(\psi'(f), p)$ and maximize it. We do it below for an important optical network topology – the unidirectional ring. Note that the difference in path capacity between the current state and the new state is either 0 or 1 for any path. Informally stated, the $M\Sigma$ algorithm chooses that wavelength that minimizes the number of paths whose capacities decrease by 1.

6.3.1 The Unidirectional Ring Network

Consider a unidirectional ring network and assume that a call arrives. We assume that a call from a node to itself may arrive, but the implementation can be easily modified if this were not the case. Assume that the links of the ring are oriented in the clockwise direction and are numbered $1, 2, \dots, N$, in the clockwise direction. Without loss of generality, let the call request use links 1 through n , and let p denote the corresponding path. For each wavelength $f \in \Omega(\psi, p)$, a cost function $\chi(p, f)$ that counts the number of paths whose capacities are decreased by 1 is computed and the wavelength f^* that minimizes the function is chosen to route the call. Let $u(f, \psi)$ be the number of fibers in the network on which wavelength f is used. Before we give the cost function, let us develop some more notation.

Define $k_i \stackrel{\text{def}}{=} \min\{c(\psi, 1, f), c(\psi, 2, f), \dots, c(\psi, i, f)\}$, and $g_i \stackrel{\text{def}}{=} \min\{c(\psi, n, f), c(\psi, n-1, f), \dots, c(\psi, n-i+1, f)\}$ for $i = 1, 2, \dots, n$. In other words, k_i is the number of fibers on which f is available on the most congested link on f of the first i links of path p , and g_i is the number of fibers on which f is available on the most congested link on f of the last i links of the path.

Let $h_i^{(l)}$ be the left gap from link 1 until a link with capacity less than k_i is encountered and, similarly, let $h_i^{(r)}$ be the right gap from link n until a link with capacity less than g_i is encountered. Formally stated, let $L_i = \{l : n+1 \leq l \leq N, c(\psi, l, f) < k_i\}$ and let $R_i = \{r : n+1 \leq r \leq N, c(\psi, r, f) < g_i\}$. If $L_i = \phi$, $h_i^{(l)} \equiv N - i$; otherwise, $h_i^{(l)} \equiv N - \max L_i$. If $R_i = \phi$, $h_i^{(r)} \equiv N - n$; otherwise, $h_i^{(r)} \equiv \min R_i - n - 1$.

The cost function $\chi(p, f)$ for wavelength f can now be written as follows:

$$\chi(p, f) = \begin{cases} N^2, & n = N, \\ N^2 - (N - n)(N - n + 1)/2, & n < N, u(\psi, f) = 0, \\ \frac{(N-n)(N-n-1)}{2} + \frac{n(n+1)}{2} + \sum_{i=1}^n (h_i^{(l)} + h_i^{(r)}), & n < N, h_n^{(l)} = h_n^{(r)} = N - n, \\ h_n^{(l)} h_n^{(r)} + \frac{n(n+1)}{2} + \sum_{i=1}^n (h_i^{(l)} + h_i^{(r)}), & n < N; h_n^{(r)}, h_n^{(l)} \neq N - n. \end{cases}$$

Similar formulas can be obtained for bidirectional ring networks too. The worst-case time complexity of the $M\Sigma$ algorithm with this implementation is $O(NF)$.

6.3.2 Arbitrary Topologies

We now discuss an implementation of the $M\Sigma$ algorithm in arbitrary mesh networks with fixed routing. Suppose we are given an arbitrary mesh network represented by a graph $G = (V, E)$ where V is the set of nodes and E is the set of links (each consisting of many fibers). Let a set of fixed paths P between node-pairs also be given. We generate a *conflict graph* $G'(V', E')$ where each $p' \in V'$ corresponds to a path $p \in P$ and an undirected edge $(p', q') \in E'$ iff the paths p and $q \in P$ share a common link.

Suppose a call is to be established on path p . For each wavelength, a cost function $\chi(p, f)$ that counts the number of paths whose capacities are decreased by one in the new state is computed as follows. First, define

$$M(r) \stackrel{\text{def}}{=} \{l \in L(r) : c(\psi, l, f) \leq c(\psi, m, f) \forall m \in L(r)\}$$

to be the set of links of path r which have the minimum capacity on wavelength f . Recall that $L(r)$ is the set of links on path r . Now, let us define

$$R_p \stackrel{\text{def}}{=} \{r' : (p', r') \in E', M(r) \cap L(p) \neq \emptyset\}$$

to be the set of those paths that have a minimum capacity link lying on p . We can now write $\chi(p, f) = |R_p|$ as the cost of using wavelength f .

Though the cost function defined above is for the $M\Sigma$ algorithm, it can be trivially modified for other algorithms such as $MW\Sigma$ and Max_Min. The worst-case time complexity of the algorithms using the conflict graph framework is $O(|P|NF)$ assuming the conflict graph is generated offline.

The implementation on arbitrary topologies can be easily extended to consider adaptive routing on a set of k fixed alternate paths. The conflict graph is generated offline by considering all the k paths for each node-pair. The algorithm assigns a cost function $\chi(p_i, f)$ for each of the k paths $p_1, p_2, \dots, p_k$ as above. The wavelength and route that minimize the cost function over all wavelengths and k paths are then chosen to establish the call. We next present an extension to multi-fiber networks of a performance model [22, 27] for single-fiber networks with and without wavelength conversion using random wavelength assignment.

6.4 An Analytical Blocking Model for Dense Multi-fiber Networks

Analytical models predicting the blocking probability (P_b) of single-fiber networks with random wavelength assignment under Poisson traffic have been presented in [18, 19, 20, 22] and in Chapter 3. The only model for multi-fiber networks was presented in [21, 25]. Let d be the number of fibers per link (link dilation degree) and let F be the number of wavelengths per fiber. The model of [21, 25] makes the following incorrect prediction: For fixed F , P_b without conversion remains *constant* as d increases and the load is equal to the value that would give a constant P_b with conversion. As simulations show in the next section, P_b without conversion approaches P_b with conversion rapidly as d increases in a mesh-torus network with Poisson traffic. We now present a straightforward extension of the model presented in [22] and Chapter 3 to reasonably dense multi-fiber networks in which link-load independence can be assumed. This model makes correct predictions in the above scenario and is numerically accurate enough to let qualitative conclusions to be drawn about the

benefits of conversion in multi-fiber networks.

Assuming uniform Poisson traffic, the offered load per link ρ can be determined given the load per wavelength per fiber δ as

$$\rho = \frac{\delta F d \bar{H}}{L}$$

where $\bar{H}$ is the average hop length and L is the number of links in the network. We ignore the effect of the reduced link load due to blocking in computing P_b as in Chapters 3 and 4. Let a wavelength be free on a link if it is available on one of the fibers on the link. The steady-state probability that y wavelengths are free on a link given m calls are active on the link can be written as

$$S(y|m) = \frac{\binom{F}{y} f((m - (F - y)d), y, d)}{\binom{Fd}{m}}. \quad (6.1)$$

Here, $f(r, s, d)$ is the number of ways of distributing r balls (one in each slot) into s sets of d slots each such that no set has balls in all d slots. Suppose $j = \lfloor \frac{r}{d} \rfloor$. $f(r, s, d)$ is given by

$$f(r, s, d) = \begin{cases} 0, & r > (d - 1)s, \quad r < 0, \\ \binom{sd}{r}, & j = 0, \\ \binom{sd}{r} - \sum_{i=1}^j \binom{s}{i} f(r - id, s - i, d), & \text{otherwise.} \end{cases}$$

The link busy wavelength distribution p'_y can be obtained from (6.1) by using an $M/M/Fd/Fd$ queueing model to compute the probability that m calls are active on a link. We thus have

$$p'_y = \frac{1}{\sum_{i=0}^{Fd} \frac{\rho^i}{i!}} \sum_{m=0}^{Fd} S(F - y|m) \frac{\rho^m}{m!}. \quad (6.2)$$

(6.2) together with (4.10)-(4.14) of Sections 4.3.2 and 4.3.3 can be used to obtain the blocking performance of a multi-fiber network topology. We apply this model to the mesh-torus network in the next section and observe that the model yields reasonably accurate results.

6.5 Results

We now present results comparing the performance of the $M\Sigma$ algorithm with that of other proposed heuristics. The comparison is done by extensive simulations assuming uniform Poisson traffic. We assume that the link dilation degree d is the same for all links and the number of wavelengths on all fibers is F .

Suppose a call on path p needs to be assigned a wavelength and the network is in state ψ . The heuristics that we simulated are: (1) FF - the First-Fit heuristic that assigns the available wavelength with lowest index [22, 32, 50], (2) MU - the Most-Used heuristic that assigns the wavelength that is used on the most number of fibers in the network [32, 50, 51, 52], (3) MP - the Min-Product heuristic that chooses the wavelength to $\min_{f \in \Omega(\psi, p)} \prod_{l \in L(p)} \{d - c(\psi, l, f)\}$ [21], (4) LL - the Least Loaded heuristic that chooses f such that $\min_{l \in L(p)} c(\psi, l, f)$ is maximum [48], (5) R - the Random wavelength assignment heuristic where one of the available wavelengths is chosen randomly, and (6) MS - the $M\Sigma$ heuristic². We also obtain the blocking performance of the network with wavelength conversion (WC). For MP, LL, and MS, ties between wavelengths are broken by picking the MU wavelength among those. This gave better results than using FF as the tie-breaker. For single-fiber networks, MP and LL reduce to the MU heuristic (tie-breaker).

In Figures 6.1(a) and 6.1(b), we plot the performance of the various algorithms against F for $d = 1$ and $d = 10$, respectively, for a unidirectional ring network. We see that in a single-fiber ring network, our algorithm performs similar to MU. Since the conflict graph of a ring network has diameter 2, the $M\Sigma$ algorithm chooses the MU wavelength with high probability when $d = 1$. When $d = 10$, our algorithm outperforms LL, the best algorithm previously proposed, sometimes by an order of magnitude. This trend in performance continues to hold as d increases to a few tens

² Performance results of the algorithm in [45] for blocking networks have not been published. We do not simulate the algorithm because it is applicable only for specific values of F .

but the difference in performance between the heuristics and WC gradually diminishes as d increases to very high values (> 100).

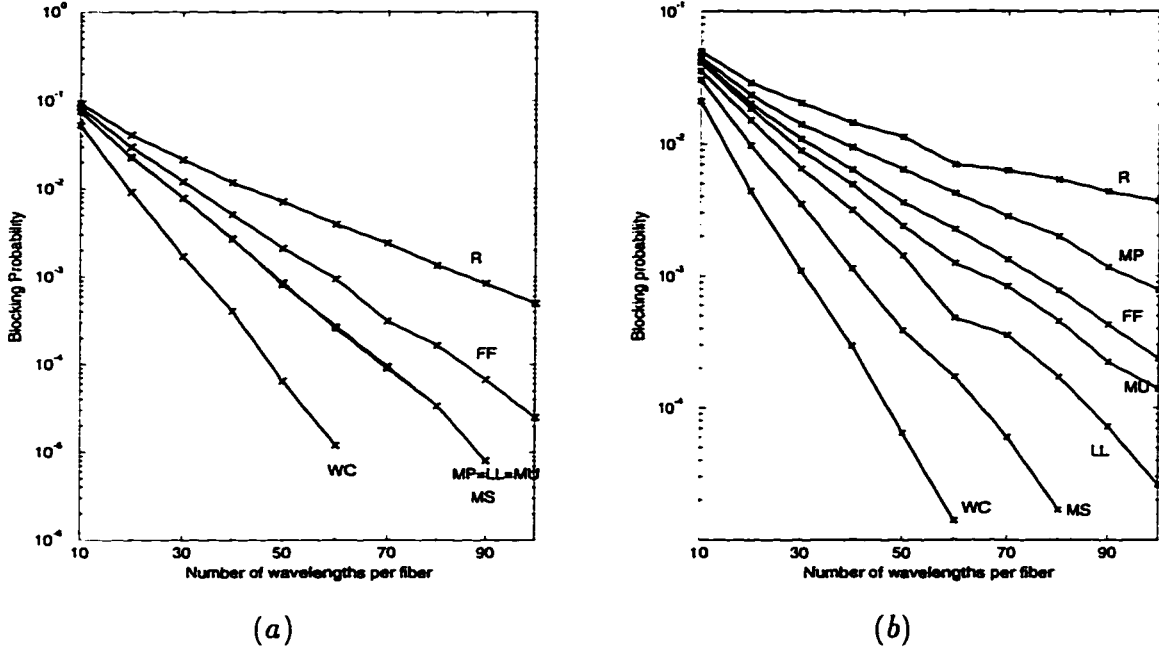
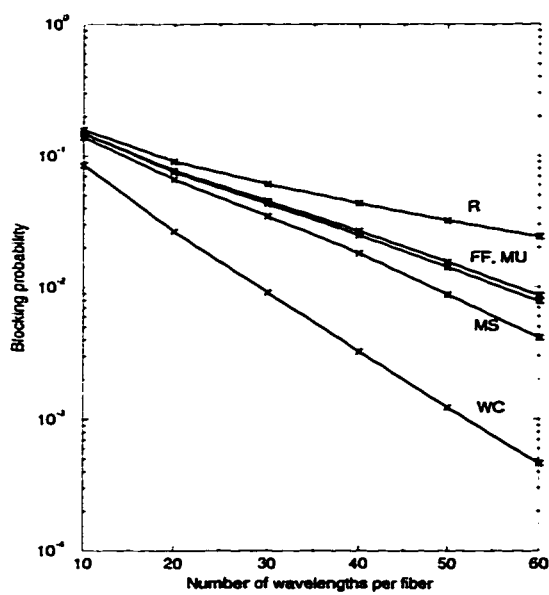
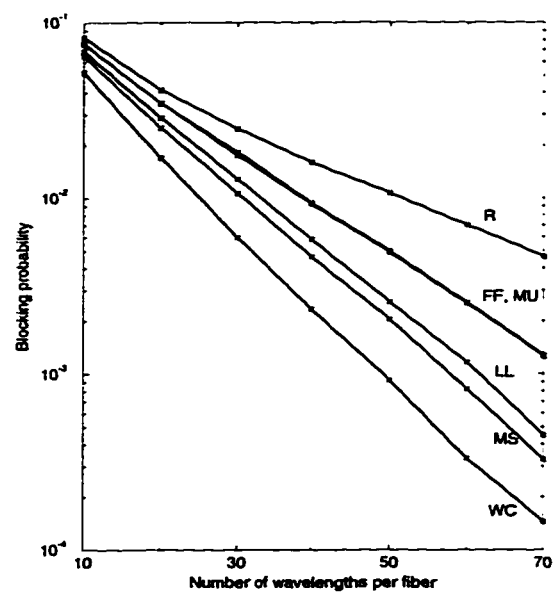


Figure 6.1: P_b vs. F for the different algorithms in a 20-node unidirectional ring (a) $d = 1$, $\text{load}/f/d = 1$ Erlang (b) $d = 10$, $\text{load}/f/d = 1.6$ Erlangs.

We next plot P_b against F for $d = 1$ and $d = 3$ for a 5×5 bidirectional mesh-torus network in Figure 6.2. Shortest-path deterministic X-Y routing [22] is assumed. In this case, our algorithm performs somewhat better than MU even when $d = 1$. The difference is however small. The performance improvement that we obtained in the ring network with increasing d is not seen here as d begins to increase. In fact, as d assumes even modest values (such as 10), the performance of all the heuristics (including R) are close to WC, suggesting that the choice of heuristic and wavelength conversion are not very important in a multi-fiber mesh-torus network. This is in contrast to single-fiber mesh-torus networks for which it has been shown in Chapter 3 and [22] that R performs poorly when compared to WC.



(a)



(b)

Figure 6.2: P_b vs. F for the different algorithms in a 5×5 bidirectional mesh-torus (a) $d = 1$, load/f/d = 25 Erlangs (b) $d = 3$, load/f/d = 31 Erlangs.

In order to study the effect of dilation degree on the performance of the various heuristics, we now plot the blocking performance against d for a fixed value of F . First, define the blocking probability loss (L_p) of a heuristic H as the ratio of P_b using H to P_b with WC for a given load.

L_p is plotted against d for the various heuristics by adjusting the load so that $P_b(WC) \approx 10^{-4}$ for a mesh-torus in Figure 6.3. The loss quickly approaches 1 for all the heuristics including R. The R model predicts this behavior to an approximation. At the moment, we do not have an explanation for why the model overestimates for $d = 1$ and underestimates for $d > 1$. The point here is that for the multi-fiber mesh-torus, random wavelength assignment is almost as good as WC.

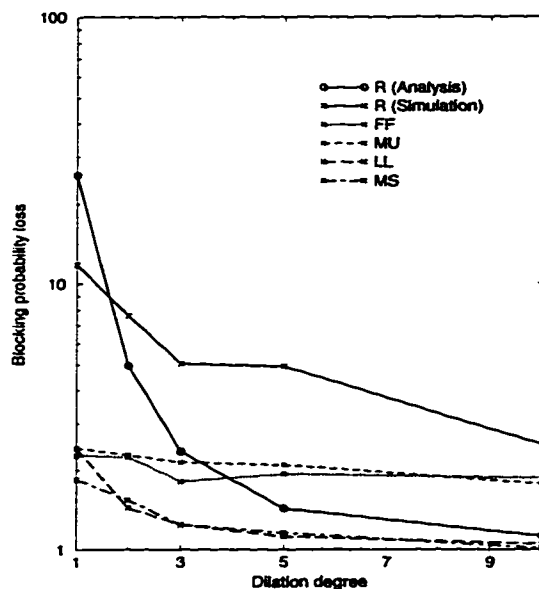


Figure 6.3: L_p vs. d for a 5×5 mesh-torus. $P_b(WC) \approx 0.0001$, $F = 10$.

On the other hand, as shown in Figure 6.4, the behavior of L_p is completely different in the case of the ring³. L_p increases with d initially and then decreases

³ Note that L_p could go below 1 since a heuristic may admit more short-path calls and block a few longer-path calls that may be admitted by WC, thereby decreasing the overall P_b of the heuristic.

very gradually for R, FF, and MU. The same is true for the other heuristics except that the increase in L_p is much lower (note the log scale) and the subsequent drop is sharper, implying that smart heuristics really pay off in multi-fiber rings with moderate dilation degrees.

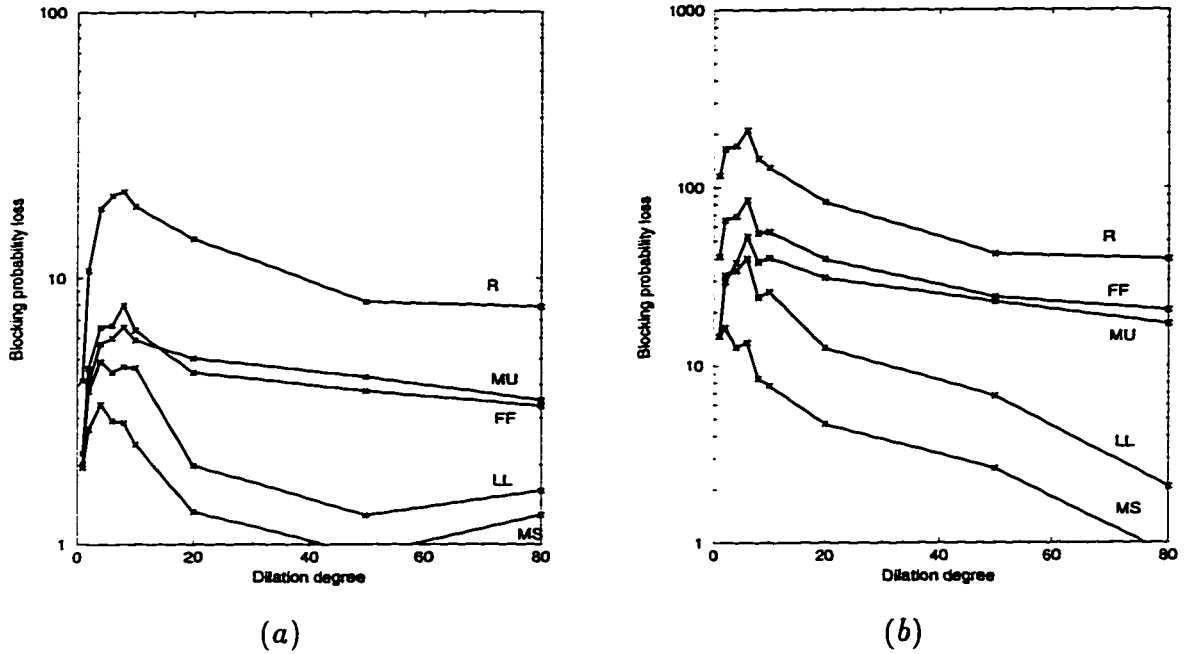


Figure 6.4: L_p vs. d for a 20-node ring. $P_b(WC) \approx 0.0001$ (a) $F = 10$ (b) $F = 60$.

In single-fiber networks, it was shown in [18, 22] and Chapter 3 that wavelength conversion is more useful in the mesh-torus (high L_p for random wavelength assignment) than in the ring. We see here that the situation in multi-fiber networks is quite the *opposite*. The decrease in L_p with d in the mesh-torus is somewhat intuitive. The behavior of L_p in the multi-fiber ring is very surprising and merits further attention.

6.6 Conclusions

In this chapter, we presented a new centralized dynamic wavelength assignment algorithm for wavelength-routed WDM networks. We also presented simple implementations of our algorithm for ring networks and for arbitrary topologies. The algorithm performs better than previously proposed algorithms in all studied cases, and the performance improvement is significant in rings with moderately large dilation degrees (5 – 40). Our algorithm's performance is better because it considers the state of the whole network after wavelength assignment and not just the path of the call. Alternate routing remains to be implemented and the performance needs to be studied. Some of the other algorithms that we have proposed may be more appropriate in certain situations, e.g., $MW\Sigma$ with weights equal to path loads may be better for non-uniform traffic. We also extended our analytical performance model [27] to the case of multi-fiber networks and showed that the model is accurate enough to make qualitative predictions for reasonably dense topologies. From our studies, we also conclude that dilation affects the performance in rings and mesh-tori in very different ways.

Chapter 7

CONCLUSIONS AND FUTURE WORK

The primary focus of this dissertation was the impact of wavelength conversion on the performance of wavelength-routed WDM networks. WDM networks employing wavelength routing is currently the most attractive option for a wide-area backbone network built on a national or international scale. Wavelength conversion has been recognized as a key capability for improving the network performance.

All-optical wavelength converters remain expensive, hard-to-design devices and a thorough understanding of their benefits is necessary for the design of future all-optical networks. In Chapter 3, we proposed the concept of sparse wavelength conversion. Networks with sparse wavelength conversion do not necessarily have wavelength conversion capability at every node. We developed analytical models for Poisson traffic to evaluate the blocking performance of such networks employing random wavelength assignment. Our analytical model captures the important effects of load correlation without increasing the computational complexity dramatically. To avoid the tedious task of computing the blocking probabilities for every possible placement of a given number converters in a network, we obtained the ensemble average performance of the network over all possible numbers of converters with a given expected value and all possible placements. This allowed a single parameter characterization of networks with sparse wavelength conversion. The results suggest that wavelength conversion does not reduce the blocking probabilities significantly in very sparsely connected networks such as the ring and very densely connected networks such as the hypercube, while yielding significant performance benefits in moderately connected networks such as the mesh-torus for uniform Poisson traffic. Furthermore, our results

indicate that typically a small conversion density (fraction of network nodes with conversion capability) may be enough to boost the performance to a certain required level.

The exact nature of traffic that future optical networks are expected to carry is not known. In Chapter 4, we attempted to study the sensitivity of the conclusions reached in Chapter 3 to the assumption of Poisson input traffic. We employed a simple model called the Bernoulli-Poisson-Pascal model to study the effect of the second moment of an input traffic (peakedness) on the benefits of wavelength conversion. The results of that chapter indicate that peakedness dramatically (and negatively) affects the blocking probabilities in both networks with and without conversion. However, networks with wavelength conversion are affected more. On the other hand, the increase in load achievable with wavelength conversion at a fixed blocking performance is not affected by traffic peakedness over a large range of peakedness values.

The concept of sparse wavelength conversion naturally gives rise to the problem of optimally placing a given number of wavelength converters in a network. We addressed this problem in Chapter 5 and provided optimal solutions to minimize the blocking probability in a path of a network, bus topologies, and ring topologies. We showed that optimal converter placement is an important issue that results in significant blocking performance improvement, especially when link load correlation is insignificant. The results obtained using our dynamic programming-based solutions indicate that it is hard to estimate optimal placements *a priori*. We also studied the effect of the blocking performance model on the optimal placement decision. It appears that the decisions are largely insensitive to the performance model; however, there were instances in which the placements obtained using two different performance models were different.

In Chapter 6, we considered networks without wavelength conversion. We assumed that the routes of calls were fixed and developed an efficient heuristic called the Max.Sum algorithm for assigning wavelengths to dynamically arriving connection

requests. Our algorithm was observed to perform better than other existing heuristics for Poisson traffic using extensive simulations. The performance improvement with the Max.Sum algorithm is considerable in ring networks with a moderate number of fibers per link (5 – 40) and a large number of wavelengths (> 20). In a mesh-torus network, we observed that the wavelength assignment algorithm is not very crucial to improving the performance when the number of fibers per link is high (> 10).

The research reported in this dissertation can be extended in several different ways. The model we presented in Chapter 3 assumed that alternate routing was not allowed. However, practical networks almost always employ alternate routing and models for wavelength conversion with alternate routing need to be developed. Furthermore, models for wavelength assignment algorithms other than random are required to study the benefits of wavelength conversion in a more practical scenario. There has been some work addressing these two aspects recently in the literature, but there is no reported work yet on the impact of sparse wavelength conversion under alternate routing and other wavelength assignment algorithms. Another possible direction that could be pursued is the performance of networks whose nodes have a limited number of transceivers. We assumed in our models that any connection can be established as long as there is a wavelength available on the fibers along the connection's path.

The evaluation of blocking performance is only one measure of analyzing the benefits of wavelength conversion. Other approaches such as capacity planning could be used. The recent literature does report some work in this area for networks with sparse wavelength conversion, but more is required for a clearer understanding.

We used a simple model to study the effect of traffic peakedness on wavelength conversion benefits. Possible future work could include other models and the effect of higher moments of traffic.

Converter placement in arbitrary topologies remains an open problem. In light of the difficult nature of the problem, a heuristic approach to place converters would

appear to be necessary. Effective heuristics, possibly based on solutions we have presented, should be developed and the importance of near-optimal placement in arbitrary topologies should be studied.

The wavelength assignment algorithm presented in Chapter 6 could be extended to assign wavelengths to a set of static requests in an attempt to minimize the number of wavelengths used. A possible way of ordering the requests for wavelength assignment is to assign wavelengths for the longest-hop calls first. Some of the other algorithms suggested in Chapter 6 such as `Max_Min` and `Max_Weighted_Sum` may be more appropriate in certain situations and results obtained using those algorithms could be interesting. We have evaluated the performance of our algorithm only on rings and mesh-tori. Performance in other topologies remains to be studied.

Optical networks are expected to play an important role in the future information infrastructure, such as the next generation Internet architecture. It is hoped that this dissertation has addressed some issues that might be useful in the design of such networks. Our development of accurate performance models and of the notion of sparse wavelength conversion may be beneficial in this context.

BIBLIOGRAPHY

- [1] P. E. Green, Jr. Optical networking update. *IEEE Journal on Selected Areas in Communications*, 14(5):764–779, June 1996.
- [2] R. Ramaswami. Multiwavelength lightwave networks for computer communication. *IEEE Communications Magazine*, 31(2):78–88, Feb. 1993.
- [3] P. E. Green, Jr. *Fiber Optic Networks*. Prentice Hall, 1992.
- [4] B. Mukherjee. WDM-based local lightwave networks – Part II: Multihop systems. *IEEE Network*, 6(4):20–32, July 1992.
- [5] B. Mukherjee. WDM-based local lightwave networks – Part I: Single-hop systems. *IEEE Network*, 6(3):12–27, May 1992.
- [6] M. S. Goodman. Multiwavelength networks and new approaches to packet switching. *IEEE Communications Magazine*, 27(10):27–35, Oct. 1989.
- [7] R. A. Barry and P. A. Humblet. Latin routers, design and implementation. *Journal of Lightwave Technology*, 11(5/6):891–899, May/June 1993.
- [8] R. A. Barry and P. A. Humblet. On the number of wavelengths and switches in all-optical networks. *IEEE Transactions on Communications*, 42(2/3/4):583–591, Feb./March/April 1994.
- [9] R. K. Pankaj and R. G. Gallager. Wavelength requirements of all-optical networks. *IEEE/ACM Transactions on Networking*, 3(3):269–280, June 1995.

- [10] R. Ramaswami and K. N. Sivarajan. Routing and wavelength assignment in all-optical networks. *IEEE/ACM Transactions on Networking*, 3(5):489–500, Oct. 1995.
- [11] T. Shiragaki, M. Fujiwara, S. Suzuki, C. Burke, and T. Shiozawa. Optical digital cross-connect system using photonic switch matrices and optical amplifiers. *Journal of Lightwave Technology*, 12(8):1490–1496, Aug. 1994.
- [12] J. Zhou, N. Park, K. J. Vahala, M. A. Newkirk, and B. I. Miller. Broadband wavelength conversion with amplification by four-wave mixing in semiconductor travelling-wave amplifiers. *Electronics Letters*, 30(11):859–860, May 1994.
- [13] C. Joergensen, S. L. Danielsen, M. Vaa, B. Mikkelsen, K. E. Stubkjaer, P. Doussiere, F. Pommerau, L. Goldstein, and M. Goix. 40 Gbit/s all-optical wavelength conversion by semiconductor optical amplifiers. *Electronics Letters*, 32(4):367–368, May 1996.
- [14] B. Mikkelsen, T. Durhuus, C. Joergensen, R. J. S. Pedersen, C. Braagaard, and K. E. Stubkjaer. Polarisation insensitive wavelength conversion of 10 Gbit/s signals with SOAs in a Michelson interferometer. *Electronics Letters*, 30(3):260–261, May 1994.
- [15] K. C. Lee and V. O. K. Li. A wavelength-convertible optical network. *Journal of Lightwave Technology*, 11(5/6):962–970, May/June 1995.
- [16] R. Ramaswami and G. H. Sasaki. Multiwavelength optical networks with limited wavelength conversion. In *Proc. INFOCOM '97*, April 1997.
- [17] O. Gerstel, R. Ramaswami, and G. H. Sasaki. Benefits of limited wavelength conversion in WDM ring networks. In *Proc. OFC '97*, Feb. 1997.

- [18] R. A. Barry and P. A. Humblet. Models of blocking probability in all-optical networks with and without wavelength changers. *IEEE Journal on Selected Areas in Communications*, 14(5):858–867, June 1996.
- [19] A. Birman. Computing approximate blocking probabilities for a class of all-optical networks. *IEEE Journal on Selected Areas in Communications*, 14(5):852–857, June 1996.
- [20] R. A. Barry and D. Marquis. Evaluation of a model of blocking probability in all-optical mesh networks without wavelength changers. In *Proc. SPIE*, pages 154–167, Oct. 1995.
- [21] G. Jeong and E. Ayanoglu. Comparison of wavelength-interchanging and wavelength-selective cross-connects in multiwavelength all-optical networks. In *Proc. INFOCOM '96*, pages 156–163, March 1996.
- [22] M. Kovačević and A. S. Acampora. Benefits of wavelength translation in all-optical clear-channel networks. *IEEE Journal on Selected Areas in Communications*, 14(5):868–880, June 1996.
- [23] J. Yates, J. Lacey, D. Everitt, and M. Summerfield. Limited-range wavelength translation in all-optical networks. In *Proc. INFOCOM '96*, pages 954–961, March 1996.
- [24] I. Chlamtac, A. Ganz, and G. Karmi. Lightpath communications: An approach to high bandwidth optical WAN's. *IEEE Transactions on Communications*, 40(7):1171–1182, July 1992.
- [25] R. A. Barry and D. Marquis. An improved model of blocking probability in all-optical networks. In *LEOS 1995 Summer Topical Meeting*, Aug. 1995.

- [26] B. S. Glance, J. M. Wiesenfeld, U. Koren, and R. W. Wilson. New advances in optical components needed for FDM optical networks. *Journal of Lightwave Technology*, 11(5/6):882–890, May/June 1993.
- [27] S. Subramaniam, M. Azizoglu, and A. K. Somani. Connectivity and sparse wavelength conversion in wavelength-routing networks. In *Proc. INFOCOM '96*, pages 148–155, March 1996.
- [28] S. Subramaniam, M. Azizoglu, and A. K. Somani. All-optical networks with sparse wavelength conversion. *IEEE/ACM Transactions on Networking*, 4(4):544–557, Aug. 1996.
- [29] A. Girard. *Routing and Dimensioning in Circuit-Switched Networks*. Addison-Wesley, 1990.
- [30] S. B. Alexander, R. S. Bondurant, D. Byrne, V. W. S. Chan, S. G. Finn, R. Gallager, B. S. Glance, H. A. Haus, P. Humblet, R. Jain, I. P. Kaminow, M. Karol, R. S. Kennedy, A. Kirby, H. Q. Le, A. A. M. Saleh, B. A. Schofield, J. H. Shapiro, N. K. Shankaranarayanan, R. E. Thomas, R. C. Williamson, and R. W. Wilson. A precompetitive consortium on wide-band all-optical networks. *Journal of Lightwave Technology*, 11(5/6):714–735, May/June 1993.
- [31] H. Harai, M. Murata, and H. Miyahara. Performance of alternate routing methods in all-optical switching networks. In *Proc. INFOCOM '97*, April 1997.
- [32] A. Mokhtar and M. Azizoglu. Adaptive wavelength routing in all-optical networks. Submitted to *IEEE/ACM Transactions on Networking*, 1997.
- [33] D. Bertsekas and R. Gallager. *Data Networks*. Prentice Hall, 1992.

- [34] G. A. Corea and V. G. Kulkarni. Shortest paths in stochastic networks with arc lengths having discrete distribution. *Networks*, 23(3):175–183, May 1993.
- [35] S. K. Walley, H. H. Tan, and A. M. Viterbi. Limit distributions for the diameter and the shortest path hop count in random graphs with positive integer edge costs. In *Proc. INFOCOM '94*, pages 618–627, June 1994.
- [36] C. Rose. Mean internodal distance in regular and random multihop networks. *IEEE Transactions on Communications*, 40(8):1310–1318, Aug. 1992.
- [37] I. P. Kaminow, C. R. Doerr, C. Dragone, T. Koch, U. Koren, A. A. M. Saleh, A. J. Kirby, C. M. Özveren, B. Schofield, R. E. Thomas, R. A. Barry, D. M. Castagnozzi, V. W. S. Chan, B. R. Hemenway, Jr., D. Marquis, S. A. Parikh, M. L. Stevens, E. A. Swanson, S. G. Finn, and R. G. Gallager. A wideband all-optical WDM network. *IEEE Journal on Selected Areas in Communications*, 14(5):780–799, June 1996.
- [38] E. A. van Doorn. Some aspects of the peakedness concept in teletraffic theory. *Elektronische Informationsverarbeitung und Kybernetik*, 22(2–3):93–104, 1986.
- [39] L. E. N. Delbrouck. A unified approximate evaluation of congestion functions for smooth and peaky traffics. *IEEE Transactions on Communications*, 29(2):85–91, Feb. 1981.
- [40] R. I. Wilkinson. Theories for toll traffic engineering in the U.S.A. *Bell Systems Technical Journal*, 35:421–514, March 1956.
- [41] A. Kuczura and D. Bajaj. A method of moments for the analysis of a switched communication network's performance. *IEEE Transactions on Communications*, 25(2):185–193, Feb. 1977.

- [42] H. Heffes and D. Lucantoni. A Markov-modulated characterization of packetized voice and data traffic and related statistical multiplexer performance. *IEEE Journal on Selected Areas in Communications*, 4(6):856–868, Sep. 1986.
- [43] M. Gondran and M. Minoux. *Graphs and Algorithms*. Wiley, 1986.
- [44] B. Narahari, S. Shende, and R. Simha. Efficient algorithms for erasure node placement on slotted dual bus networks. *IEEE/ACM Transactions on Networking*, 4(5):779–784, Oct. 1996.
- [45] O. Gerstel and S. Kutten. Dynamic wavelength allocation in all-optical ring networks. In *Proc. ICC '97*, June 1997.
- [46] Z. Zhang and A. S. Acampora. A heuristic wavelength assignment algorithm for multihop WDM networks with wavelength routing and wavelength re-use. *IEEE/ACM Transactions on Networking*, 3(3):281–288, June 1995.
- [47] N. Wauters and P. Demeester. Design of the optical path layer in multiwavelength cross-connected networks. *IEEE Journal on Selected Areas in Communications*, 14(5):881–892, June 1996.
- [48] E. Karasan and E. Ayanoglu. Effects of wavelength routing and selection algorithms on wavelength conversion gain in WDM optical networks. In *LEOS 1996 Summer Topical Meeting on Broadband Optical Networks*, Aug. 1996.
- [49] R. A. Barry and S. Subramaniam. The MAX-SUM wavelength assignment algorithm for WDM ring networks. In *Proc. OFC '97*, June 1997.
- [50] A. Mokhtar. *Switching, routing, and multiaccess in all-optical networks*. PhD thesis, University of Washington, Seattle, 1997.

- [51] I. Chlamtac, A. Ganz, and G. Karmi. Purely optical networks for terabit communication. In *Proc. INFOCOM '89*, pages 887–896, April 1989.
- [52] K. Bala, T. E. Stern, D. Simchi-Levi, and K. Bala. Routing in linear lightwave networks. *IEEE/ACM Transactions on Networking*, 3(4):459–469, Aug. 1995.

Appendix A

AN ILP FOR CONVERTER PLACEMENT

This appendix provides an integer-linear programming formulation for a version of the converter placement problem with static traffic. For convenience, let P denote the number of paths (or (s, d) pairs) and λ_i denote the number of sessions to be established on path i . In our formulation here, we assume the routing to be given. An alternative formulation in which the routing is not given can be obtained in a similar manner. Let N denote the number of nodes in the network, L the number of links, and F the number of wavelengths per fiber. The number of fibers per link is assumed to be 1. Let M_i denote the number of sessions that are actually established on path i for any converter placement and wavelength assignment (WA) algorithm.

Let $\mathbf{B} = (b_{ij})$ denote the path-link incidence matrix, i.e.,

$$b_{ij} = \begin{cases} 1, & \text{if link } j \text{ is on path } i, \\ 0, & \text{otherwise.} \end{cases}$$

Let $\mathbf{D} = (d_{ij})$ denote the link-node incidence matrix, i.e.,

$$d_{ij} = \begin{cases} 1, & \text{if link } i \text{ is incident to node } j, \\ 0, & \text{otherwise.} \end{cases}$$

The allocation of converters and the corresponding WA algorithm consistent with that placement can then be described by an N -vector $\mathbf{A} = (a_i)$ and a $P \times L \times F$ matrix $\mathbf{C} = (c_{ijk})$, where

$$a_i = \begin{cases} 1, & \text{if a converter is placed at node } i, \\ 0, & \text{otherwise,} \end{cases}$$

denotes the converter location vector, and

$$c_{ijk} = \begin{cases} 1, & \text{if wavelength } k \text{ is assigned to link } j \text{ on path } i, \\ 0, & \text{otherwise,} \end{cases}$$

represents the allocation of wavelengths to links on paths.

Then, the optimal converter placement and the corresponding RWA algorithm for the deterministic case can be obtained by solving the following ILP.

Maximize the carried traffic $\sum_{i=1}^P M_i$ subject to the following constraints.

1. $M_i \geq 0$, integer, $i = 1, \dots, N$,
2. $0 \leq a_i \leq 1$, integer, $i = 1, \dots, N$,
3. $c_{ijk} \geq 0$, integer, $i = 1, \dots, P$, $j = 1, \dots, L$, $k = 1, \dots, F$,
4. $\sum_{i=1}^P c_{ijk} \leq 1 \quad \forall j, k$ (a wavelength k can be used on a link j at most once),
5. $M_i \leq \sum_{k=1}^F c_{ijk}, \quad \forall i, j$ (the number of connections on a path is $\leq$ the number of wavelengths assigned on any link of the path),
6. $M_i \leq \lambda_i$ (traffic demands),
7. $\sum_{k=1}^F c_{ijk} \leq b_{ij} \quad \forall i, j$ (a wavelength can be assigned to a link j on path i only if link j is used on path i),
8. $c_{ijk} - c_{ilk} + (1 - a_n)d_{jn}d_{ln}b_{ij}b_{il} \leq 1 \quad \forall i, j, k, l, n$ (if link j is incident to node n and link l is incident to node n , and if there is no converter at node n , and if both links j and l lie on the same path i , then wavelength k should be assigned to both j and l or to neither of them), and
9. $\sum_{i=1}^N a_i \leq K$ (constraint on the number of converters).

VITA

Suresh Subramaniam was born in Bhadravati, India on March 14, 1968. He received the B.E. degree in Electronics and Communication Engineering from Anna University (Guindy), Madras, India, in 1988, the M.S.E.E. degree from Tulane University, New Orleans, LA, in 1993 and the Ph.D. degree in Electrical Engineering from the University of Washington, Seattle, WA, in 1997. He was employed as an R&D hardware engineer in HCL Ltd., Madras, India, from 1988 to 1990, and was a graduate student in the Department of Electrical Communication Engineering, Indian Institute of Science, Bangalore, in 1990-91.

From September 1997, he will be an Assistant Professor in the Department of Electrical Engineering and Computer Science at The George Washington University, Washington, DC.