



uOttawa

L'Université canadienne
Canada's university

**FACULTÉ DES ÉTUDES SUPÉRIEURES
ET POSTDOCTORALES**



**FACULTY OF GRADUATE AND
POSTDOCTORAL STUDIES**

Qi Guo

AUTEUR DE LA THÈSE / AUTHOR OF THESIS

M.A.Sc. (Electrical Engineering)

GRADE / DEGREE

School of Information Technology and Engineering

FACULTÉ, ÉCOLE, DÉPARTEMENT / FACULTY, SCHOOL, DEPARTMENT

**Availability-Constrained Shared Backup Path Protection for GMPLS-Based Spare Capacity
Reconfiguration**

TITRE DE LA THÈSE / TITLE OF THESIS

H. Mouftah

DIRECTEUR (DIRECTRICE) DE LA THÈSE / THESIS SUPERVISOR

CO-DIRECTEUR (CO-DIRECTRICE) DE LA THÈSE / THESIS CO-SUPERVISOR

EXAMINATEURS (EXAMINATRICES) DE LA THÈSE / THESIS EXAMINERS

Peter X. Liu

A. Boukerche

Gary W. Slater

Le Doyen de la Faculté des études supérieures et postdoctorales / Dean of the Faculty of Graduate and Postdoctoral Studies

Availability-Constrained Shared Backup Path Protection for GMPLS-Based Spare Capacity Reconfiguration

by

Qi Guo

A thesis submitted to the
Faculty of Graduate and Postdoctoral Studies

In partial fulfillment of the requirements

For the M.A.Sc. degree in
Electrical and Computer Engineering

School of Information Technology and Engineering

Faculty of Engineering

University of Ottawa

© Qi Guo, Ottawa, Canada, 2006



Library and
Archives Canada

Bibliothèque et
Archives Canada

Published Heritage
Branch

Direction du
Patrimoine de l'édition

395 Wellington Street
Ottawa ON K1A 0N4
Canada

395, rue Wellington
Ottawa ON K1A 0N4
Canada

Your file Votre référence

ISBN: 978-0-494-25781-4

Our file Notre référence

ISBN: 978-0-494-25781-4

NOTICE:

The author has granted a non-exclusive license allowing Library and Archives Canada to reproduce, publish, archive, preserve, conserve, communicate to the public by telecommunication or on the Internet, loan, distribute and sell theses worldwide, for commercial or non-commercial purposes, in microform, paper, electronic and/or any other formats.

The author retains copyright ownership and moral rights in this thesis. Neither the thesis nor substantial extracts from it may be printed or otherwise reproduced without the author's permission.

AVIS:

L'auteur a accordé une licence non exclusive permettant à la Bibliothèque et Archives Canada de reproduire, publier, archiver, sauvegarder, conserver, transmettre au public par télécommunication ou par l'Internet, prêter, distribuer et vendre des thèses partout dans le monde, à des fins commerciales ou autres, sur support microforme, papier, électronique et/ou autres formats.

L'auteur conserve la propriété du droit d'auteur et des droits moraux qui protègent cette thèse. Ni la thèse ni des extraits substantiels de celle-ci ne doivent être imprimés ou autrement reproduits sans son autorisation.

In compliance with the Canadian Privacy Act some supporting forms may have been removed from this thesis.

Conformément à la loi canadienne sur la protection de la vie privée, quelques formulaires secondaires ont été enlevés de cette thèse.

While these forms may be included in the document page count, their removal does not represent any loss of content from the thesis.

Bien que ces formulaires aient inclus dans la pagination, il n'y aura aucun contenu manquant.


Canada

Abstract

Shared Backup Path Protection (SBPP) has been widely studied in the Generalized Multi-Protocol Label Switching (GMPLS) networks due to its efficient spare capacity sharing and flexibility in service provisioning. This thesis presents two policy-based models for evaluating the end-to-end (E2E) availability of an SBPP connection by assuming that no more than two simultaneous failures could possibly occur in the network. To minimize the redundancy while meeting the E2E availability requirement, a new parameter is defined for each connection, called protection level, which creates a framework of partial restoration from any unexpected failure. Based on the proposed availability model, two novel policy-based Linear Programming (LP) formulations are introduced - called failure-dependent and failure-independent policies, which aim to reconfigure the spare capacity allocation for dynamic provisioning of SBPP connections.

Extensive simulations are conducted to validate the proposed availability model and demonstrate the effectiveness of the spare capacity reconfiguration architecture. The proposed availability-aware spare capacity reconfiguration (SCR) approaches are then implemented on top of a well known survivable routing scheme - Successive Survivable

Routing (SSR), where the spare capacity saving ratio is taken as the performance measure. We show that the proposed SCR framework is an effective approach for achieving the GMPLS-based recovery in packet-switched networks.

Acknowledgements

First, I would like to express my sincere gratitude to my supervisor, Professor Hussein T. Mouftah, for his outstanding guidance, encouragement, understanding, and help through the course of my graduate study. I am inspired by his boundless enthusiasm, dedication to excellence, and careful attention to detail. I feel privileged to have had the opportunity to study under him.

I would also like to thank Professor Pin-Han Ho of the Department of Electrical and Computer Engineering, University of Waterloo, for his effort in providing invaluable guidance and inspired ideas throughout this research. His passion for learning, extensive knowledge, and extraordinary ability to produce great scholarly work has motivated me. I sincerely appreciate all of his support and assistance.

To all the friends and colleagues at University of Ottawa and University of Waterloo, thanks for all of the good time and the assistance, both technical and non-technical.

Finally, I am deeply grateful to my parents and my brother, for their love, care, support and encouragement, without which I would not have come this far.

Page v missing.

Contents

Abstract	ii
Acknowledgements	iv
List of Figures	ix
List of Tables	xii
Acronyms	xiii
1 Introduction	1
1.1 Background	1
1.2 Motivation and Objectives	2
1.3 Thesis Contributions	6
1.4 Thesis Outline	7
2 Background and Literature Survey	8
2.1 Introduction	8

2.2	IP/MPLS over WDM Networks	9
2.2.1	Traditional IP Routing and Forwarding	9
2.2.2	MPLS	10
2.2.3	GMPLS	12
2.2.4	Network Models in IP/MPLS over WDM Networks	13
2.2.4.1	The Peer Model	13
2.2.4.2	The Overlay Model	15
2.2.4.3	The Augmented Model	15
2.2.4.4	Comparison of the Three Network Models	16
2.3	Survivability in Mesh Networks	17
2.3.1	Survivability Schemes	17
2.3.1.1	Protection	18
2.3.1.2	Restoration	19
2.3.1.3	Protection vs. Restoration	20
2.3.1.4	Shared Backup Path Protection	21
2.3.2	Single-layer and Multilayer Survivability	23
2.3.2.1	Recovery in the Optical Layer Only	23
2.3.2.2	Recovery in the IP/MPLS Layer Only	24
2.3.2.3	Layer Interworking	24
2.4	Connection Availability Analysis in Mesh Networks	26
2.4.1	Mathematical Definitions	26

2.4.1.1	Reliability	26
2.4.1.2	Availability	27
2.4.2	Availability Analysis	29
2.4.2.1	Network Component Availability	30
2.4.2.2	Availability of an End-to-End Path	30
2.4.2.3	Availability of Path-Protected Connections	31
2.5	Summary	35
3	Availability-Aware Spare Capacity Allocation and Reconfiguration	36
3.1	Introduction	36
3.2	Related Work	37
3.3	System Formulation	40
3.4	Problem Formulation	44
3.5	Proposed Spare Capacity Allocation Models	46
3.5.1	Failure-Independent SCA	47
3.5.1.1	Mathematical Formulation	47
3.5.1.2	Complexity	50
3.5.2	Failure-Dependent (FD) SCA	51
3.5.2.1	Mathematical Formulation	51
3.5.2.2	Complexity	52
3.6	Proposed Spare Capacity Reconfiguration Architecture	53
3.7	Summary	56

4	Simulation Results and Analysis	57
4.1	Introduction	57
4.2	Availability Model Validation	58
4.2.1	Simulation Setup	58
4.2.2	Results from Failure-Independent Availability Model	60
4.2.3	Results from Failure-Dependent Availability Model	66
4.3	Performance of Spare Capacity Reconfiguration Architecture	70
4.3.1	Simulation Setup	70
4.3.2	Numerical Results and Analysis	71
4.4	Summary	75
5	Conclusions and Future Work	77
5.1	Conclusions	77
5.2	Future Work	78
	Bibliography	80
A	Computation of the Stationary Probabilities of Failure Patterns	90
B	Confidence Interval	94

List of Figures

2.1	Network architecture with the peer model	14
2.2	Network architecture with the overlay model	14
2.3	Protection and restoration schemes in mesh networks	20
2.4	An illustration of MTTF and MTTR	29
2.5	An example of dedicated path protection	33
2.6	An example of shared backup path protection	34
3.1	Flow chart of the proposed SCR architecture	55
4.1	US network topology	59
4.2	pan-European network topology	59
4.3	Comparison of FID-based availabilities for each connection (US network)	61
4.4	Comparison of FID-based availabilities for each connection (pan-European)	61
4.5	Comparison of FID-based availabilities vs. protection level (US network)	63
4.6	Comparison of FID-based availabilities vs. protection level (pan-European)	63
4.7	Comparison of FID-based availabilities vs. 1/MTTF (US network)	65

4.8	Comparison of FID-based availabilities vs. $1/\text{MTTF}$ (pan-European)	65
4.9	Comparison of FD-based availabilities for each connection (US network)	67
4.10	Comparison of FD-based availabilities for each connection (pan-European)	67
4.11	Comparison of FD-based availabilities vs. $1/\text{MTTF}$ (US network)	68
4.12	Comparison of FD-based availabilities vs. $1/\text{MTTF}$ (pan-European)	68
4.13	Spare capacity saving ratio vs. availability requirement (US network)	73
4.14	Spare capacity saving ratio vs. availability requirement (pan-European)	73
4.15	Protection level vs. availability requirement (US network)	75
4.16	Protection level vs. availability requirement (pan-European)	75
A.1	The continuous time Markov chain.	92

List of Tables

2.1	Network Models in IP/MPLS over WDM Networks	16
4.1	Topological parameters of the studied networks	60

Acronyms

ASON	Automatic Switched Optical Network
CR-LDP	Constrained Label Distribution Protocol
E2E	End-to-End
DBPP	Dedicated Backup Path Protection
FEC	Forwarding Equivalent Class
FD	Failure-Dependent
FID	Failure-Independent
GL	Generalized Label
GMPLS	Generalized Multi-Protocol Label Switching
IETF	Internet Engineering Task Force
IGP	Interior Gateway Protocol
ILP	Integer Linear Programming
IP	Internet Protocol
IS-IS	Intermediate-System to Intermediate-System
ISP	Internet Service Provider
LSA	Link State Advertisement
LMP	Link Management Protocol
LP	Linear Programming
LSP	Label Switched Path
LSR	Label Switched Router
MPLS	Multi-Protocol Label Switching
MTBF	Mean Time Between Failure

MTTF	Mean Time To Failure
MTTR	Mean Time To Repair
NGI	Next Generation Internet
NMS	Network Management System
OSPF	Open Shortest Path First
OXC	Optical Cross Connect
QoP	Quality of Protection
QoS	Quality of Service
RSVP-TE	Resource Reservation Protocol with TE
SBPP	Shared Backup Path Protection
SCA	Spare Capacity Allocation
SCR	Spare Capacity Reconfiguration
SLA	Service Level Agreement
SRG	Shared Risk Group
SSP	Shared Segment Protection
SSR	Successive Survivable Routing
TDM	Time-Division Multiplexing
TE	Traffic Engineering
TCP	Transmission Control Protocol
UNI	User-Network Interface
VoIP	Voice over IP
WDM	Wavelength-Division Multiplexing

Chapter 1

Introduction

1.1 Background

The steady growth in the Internet on mission-critical business services and connection-oriented real-time multimedia applications such as Voice over IP (VoIP) and video streaming has addressed stringent demands for guaranteed service continuity and Quality of Service (QoS) in the backbone networks. Current IP backbone networks are moving toward a two-layer structure, where the top layer carries different communication services based on MPLS bandwidth provisioning, while the lower layer is formed by an optical transport network built with point-to-point wavelength-division multiplexing (WDM) transmission system and optical switching facilities. With such network architecture, a short period of interruption due to network hardware failures can disrupt thousands of connections and cause the loss of a huge amount of data.

Extensive research efforts have been addressed in recovery of single or multiple unexpected failures. Clearly, any protection/restoration mechanism is devised to increase the availability of the supported services. As the Internet evolves to the GMPLS paradigm where a connection-oriented environment addressing various QoS requirements is supported, the availability for each label switched path (LSP) for a specific service (e.g., VoIP, Transmission Control Protocol (TCP), or real-time multimedia streaming, etc.) is of great interest. This is also referred to as *service availability* or *end-to-end (E2E) availability* that can be taken as a critical performance metric on how well the network services can be supported and operated in an E2E sense. Service availability can be defined as the probability that the connection will be found in the operating state at a random time in the future [1]. Service availability requirement is usually decided by the customer application and stated in the service level agreement (SLA) along with revenue and penalty. Satisfying customer's service availability requirements to avoid penalty while minimizing the allocated resources is a major concern to a service provider.

1.2 Motivation and Objectives

To improve the E2E availability, it has been well proved that allocating redundant network resources is the best policy in the network layer when the physical availability of each network component is constant. The allocation of redundant network resources (also called protection) must be done for an LSP before any failure interrupting the LSP occurs, which is also referred to as *survivable routing*. In the dynamic GMPLS-based

bandwidth provisioning scenario, a working LSP could be equipped with one or multiple shared risk group (SRG)-disjoint backup LSPs (or path segments) such that an unexpected interruption on the working LSP could be automatically restored. A number of protection schemes have been proposed and extensively investigated in the past, such as shared backup path protection (SBPP) [2][3][4], dedicated backup path protection (DBPP) [5], shared segment protection (SSP) [6][7][8][9], etc. All of them have a design goal of reducing/minimizing the allocated spare capacity subject to different constraints and failure scenarios, such as a recovery time constraint, availability constraint, SRG-disjointedness constraint, or a guarantee in terms of survivability under one or two simultaneous failures, etc.

With DBPP, a dedicated SRG-disjoint backup LSP is established for the corresponding working LSP. In a typical SBPP implementation, on the other hand, an SRG-disjoint backup LSP is set up for the corresponding working LSP as in DBPP case, while the spare resource along the backup LSP can be shared by other backup LSPs whose working LSPs do not share a common failure with each other. This single backup path sharing condition ensures that all working LSPs can be fully restorable from any single failure since at most one of the working LSPs could possibly be hit by the failure.

Compared with DBPP, SBPP has been considered as a more aggressive spare capacity allocation strategy that can significantly reduce the required spare capacity by enabling spare resource sharing among different backup LSPs while yielding a similar level of E2E availability. Most of the previous studies on SBPP were conducted such that 100%

restorability for any single failure can be achieved. However, as the networks grow in size, and the probability of having two simultaneous failures is getting larger, it becomes a more interesting problem to investigate the availability impairment due to the dual-failure events on a DBPP or SBPP connection. Under such circumstance, DBPP and SBPP are subject to different extents of availability impairments.

A DBPP connection is unavailable only when a failure event interrupts both the working and backup LSPs regardless of the state of any other working traffic, while an SBPP connection could be disrupted not only by a failure event that interrupts both its working and backup LSPs, but also by a failure event that interrupts the working LSP of this connection and other connections which share backup resources along their backup LSPs. It is easy to find that the saving in the consumption of spare capacity in SBPP is achieved at the cost of slight E2E availability degradation and a more complicated spare capacity allocation process. In this thesis, we are particularly interested in the availability of connections with SBPP.

To our best knowledge, the majority of previous work in the area of availability evaluation and modeling for connections with SBPP has focused on the design in the optical layer [10][11][12][13][14][15]. In all of these previous works, the source node of a working path switches 100% of its bandwidth over to its protection path when failure occurs to the working path. This policy is necessary for the restoration of connections with indivisible bandwidth such as in the optical layer. In IP/MPLS layer, an LSP may support numerous independent service sessions such that dropping any/some of

them would not affect others. Thus, it could be unnecessary to require the interrupted working LSP to be either 100% restored or non-restorable.

On the other hand, it is envisioned that making the working LSP partially restorable would greatly improve the design flexibility, restoration granularity, and capacity efficiency in the event that the E2E availability constraint on each connection is the ultimate goal of network operation instead of whether the LSP can be 100% restored in presence of some number of simultaneous failures. Obviously, the latter design objective mentioned in the above has been the main focus of most of the reported studies in the past decade, and unfortunately, has little concern with the end user perception. Also, the network control and management system require an integrated strategy to perform availability-aware spare capacity reconfiguration in a dynamic network environment so as to optimize the resource allocation.

This thesis is committed to providing a solution to minimize the spare capacity allocation through availability-aware spare capacity reconfiguration in a GMPLS-based network. In particular, our objectives are as follows:

- To provide a mathematical model for evaluating the availability of SBPP connections with partial restoration.
- To provide mathematical formulations for spare capacity allocation considering the availability constraints of connections.
- To propose an availability-aware spare capacity reconfiguration architecture for

GMPLS-based dynamic bandwidth provisioning.

1.3 Thesis Contributions

The contributions of this thesis can be summarized as follows:

- A policy-based mathematical model for evaluating the availability of SBPP connections is proposed by highlighting the concept of *spare capacity availability* instead of simply *physical component availability*. Our model differs from existing models for availability evaluations in that it has finer granularity to provide protections.
- Two linear programming (LP) formulations are developed to perform spare capacity allocation for minimizing the cost of resource allocation while satisfying the availability requirements of SBPP connections based on the proposed availability model.
- An availability-aware spare capacity reconfiguration (SCR) architecture is proposed to dynamically provision availability-constrained SBPP connections with minimized spare capacity.
- Extensive simulations are conducted to validate the proposed availability model and demonstrate the effectiveness of the SCR architecture.

1.4 Thesis Outline

The thesis is organized as follows. Chapter 2 provides an overview of GMPLS-based backbone network architecture, survivability schemes and the theoretical analysis of service availability in mesh networks. Chapter 3 presents the proposed E2E availability model for connections with SBPP. Two LP formulations based on the availability model are developed to conduct spare capacity allocation. An availability-aware spare capacity reconfiguration architecture is also presented in this chapter. In Chapter 4, we present the simulation models, numerical results and performance analysis. Finally, Chapter 5 concludes this study and provides research intents in the future.

Chapter 2

Background and Literature Survey

2.1 Introduction

In this chapter, we provide background information on IP/MPLS over WDM networks and service availability analysis in mesh networks in order to facilitate the discussions in the following chapters. The chapter starts by introducing the control planes and network models of the two-layer backbone network architecture. Then the general survivability schemes in mesh networks are presented. Different approaches for providing recovery in multilayer networks are then explained. Finally, we present the mathematical definitions and analysis for service availability.

2.2 IP/MPLS over WDM Networks

2.2.1 Traditional IP Routing and Forwarding

As the Internet experiences a tremendous growth in the amount of users as well as new real-time services and applications in business and consumer markets each year, the IP traffic is becoming the majority of bandwidth carried by backbone transport networks.

In IP forwarding, a router forwards an IP packet based on the longest match for the packet's destination IP address in its forwarding table. As the packet traverses the network, each router in turn forwards the packet by reexamining its destination IP address. There is no end-to-end network connection between a source-destination (S-D) pair. In an IP-based network, traffic engineering (TE) is realized by simply manipulating cost functions and link-state metrics of interior gateway protocol (IGP), such as open shortest path first (OSPF) and intermediate-system to intermediate-system (IS-IS). Traffic Engineering (TE) is the process of controlling how traffic flows through one's network so as to optimize resource utilization and network performance [16]. Given the explosive growth of IP traffic, a number of limitations to bandwidth provisioning, TE requirements and QoS guarantees have emerged. For example, although survivability can be achieved by traditional IP routing algorithms that automatically reroute packets around a failure through routing table updates, it will take a substantial amount of time to recover from a failure, which can be in the order of several seconds to minutes and can cause serious disruption of service. This is unacceptable for many applications that require real-time

service.

2.2.2 MPLS

The development of Multi-Protocol Label Switching (MPLS) [17] technology in Internet Engineering Task Force (IETF) brings various desirable features to IP network such as TE capability, QoS support, and end-to-end backup LSPs in the event of network failures.

In MPLS [17], packets are encapsulated at ingress nodes with labels that are then used to forward the packets along label switched paths (LSPs). These LSPs can be thought of as virtual traffic trunks that carry flow aggregates generated by classifying the packets arriving at the edge or ingress nodes of an MPLS network into forwarding equivalent classes (FECs). The classification into FECs is done using packet filters that examine header fields such as source address, destination address, etc. The purpose of classifying packets into FECs is to enable the service provider to traffic engineer the network and route each FEC in a specified manner. Each packet is assigned a label associated with the FEC it belongs to. The labels are sent along with packets to the next hop, where a label switched router (LSR) uses a label forwarding table and the MPLS label to switch packets.

Signaling protocols extended with TE capabilities are used to distribute label information to establish and control an LSP. The two signaling protocols are the Resource Reservation Protocol with TE extension (RSVP-TE) [18] and the Constrained Label Distribution Protocol (CR-LDP) [19]. These protocols establish LSPs by either calculating

the path at the source node and explicitly routing the setup packets, or doing routing on a per-hop basis, and each router determines the next router along the path. Two routing protocols, extended with TE capabilities that are to be used in the IP/MPLS layer are OSPF-TE [20] and IS-IS-TE [21].

Recent progresses in WDM optical networks have dramatically driven the cost down and the bandwidth up, and the development of gigabit/terabit routers has made it possible to aggregate the lower data streams into streams suitable for WDM optical networks [22]. Hence, a network architecture of IP/MPLS over WDM is expected to form the base of next generation internet (NGI) backbone networks [23].

In the IP/MPLS over WDM network, the nodes in the optical layer are WDM-enabled optical cross connects (OXC) that are connected with fiber cables. Links between the nodes are formed by a number of wavelength channels. The LSRs are interconnected by intelligent optical core networks that provide point-to-point connectivity in the form of lightpaths with granularity of a whole wavelength. The resulting LSPs may traverse more than one lightpath. The logical topology seen by the IP/MPLS layer is the topology of the LSRs with logical links (or IP links).

The multilayer network architecture has led service providers to consider and plan the infrastructure integration paradigm one step further by integrating the optical layer into the MPLS control plane based on the emerging GMPLS architecture [24].

2.2.3 GMPLS

MPLS was developed exclusively for packet networks and supported mainly by routers and data switches. In contrast, GMPLS can be supported by a variety of optical platforms, such as OXCs and WDM systems [25]. GMPLS provides a common unified control plane to manage and provision different networks. It does not restrict the way the layers work together. Instead, GMPLS allows multiple layers to collaborate at the discretion of network operators.

MPLS is designed so that the control plane is logically separated from the data plane. GMPLS extends this concept to allow the control plane and data plane to be physically separated. Thus, the signaling and routing protocols used by MPLS technology have to be enhanced to make them suitable for circuit switching networks.

GMPLS extensions to RSVP-TE and CR-LDP signaling protocols are defined in [26][27][28] respectively. These changes are implied by the introduction of the generalized label (GL), and allow nodes to distribute GLs and perform configuration of nodes with different switching capabilities along an LSP. As for routing protocols, OSPF and IS-IS have been extended to allow dissemination of information relevant to the time-division multiplexing (TDM) and optical domains [29][30]. Furthermore, GMPLS introduces a new link management protocol (LMP) to address issues related to failure monitoring for out-of-band control channel and data links [31].

An important difference between GMPLS LSPs in the optical layer and MPLS LSPs in the IP/MPLS layer is that, in the former case, zero-bandwidth paths cannot be estab-

lished for later use. In the IP-layer case, MPLS LSPs may be established whereby if no packets are switched into the links along the path, no bandwidth is consumed. Switching packets onto these predefined paths is simple and rapid, while in the optical-layer case, merging of multiple circuits into a single outgoing circuit at the same bit rate is generally not possible [32]. Also, GMPLS LSPs require discrete units of bandwidth allocation, while the granularity of MPLS LSPs does not have to be discrete and finer.

For IP/MPLS over WDM network, based on how much and what kind of network information can be exchanged between IP/MPLS layer and WDM layer, three interconnection models are defined in [23]: peer model, overlay model, and augmented model. The basic features and functions of the three models are presented below.

2.2.4 Network Models in IP/MPLS over WDM Networks

2.2.4.1 The Peer Model

Figure 2.1 shows the network architecture of the peer model. Each LSR keeps information about the topology and the status of physical links (e.g. availability of each wavelength) in the WDM layer as well as IP links in the IP/MPLS layer. Entities of IP and WDM networks interact like peers. The network has single instances of control and management planes, which are common for the whole network. GMPLS targets this model of control plane architecture. Such a model may be appropriate when the transport and service networks are operated by a single entity.

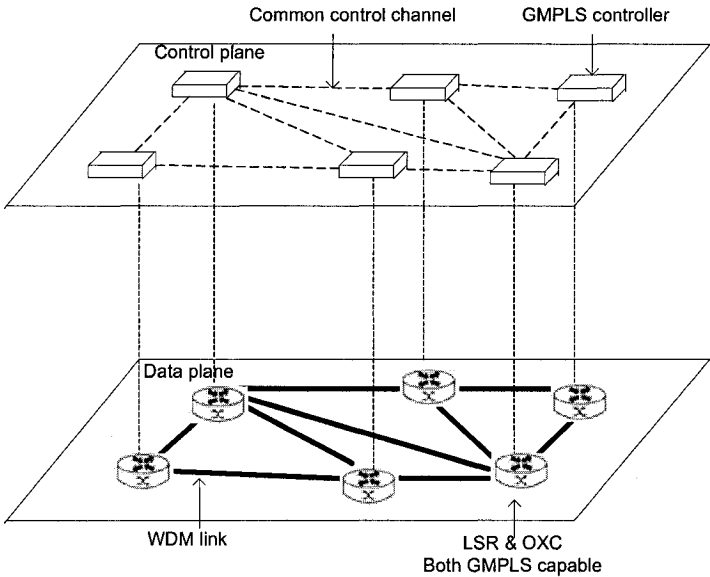


Figure 2.1: Network architecture with the peer model

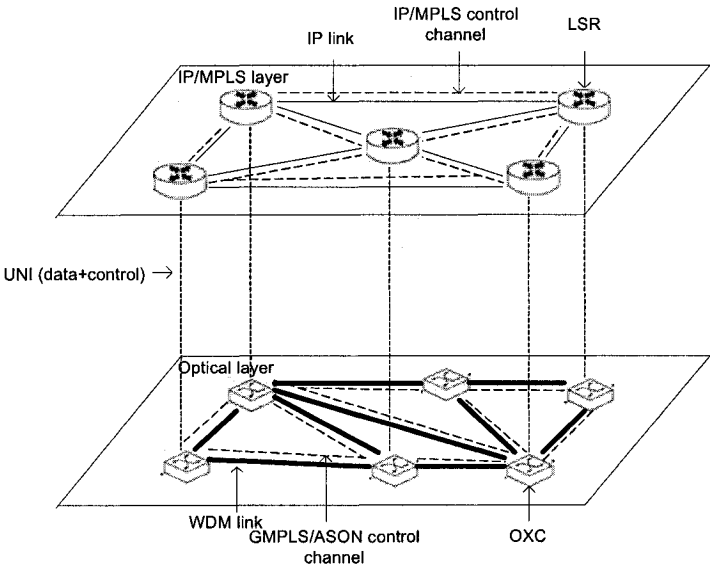


Figure 2.2: Network architecture with the overlay model

2.2.4.2 The Overlay Model

In the overlay model, a network is seen as composed of two independent layers of resources: IP/MPLS layer and WDM optical layer which have a client-server relationship as shown in Figure 2.2. There is no specific network information exchanged between the layers. Both layers have their own transport, control and management planes. Communication between the layers uses a User-Network Interface (UNI) [33]. UNI defines an interface between layers such that two adjacent layers can exchange service requests and responses as well as summarized network information if necessary.

The Automatic Switched Optical Network (ASON) [34] model defined by ITU-T targets the overlay model. In ASON model, the optical network has the ability to provision lightpaths automatically (i.e., without intrusion of management system) on demand from the client layer. Since no protocols have been developed by ITU-T to implement the intelligence of ASON, ITU-T is closely collaborating with IETF to adapt some GMPLS protocols. The GMPLS suite of protocols is expected to support new capabilities and functionalities for ASON [35].

2.2.4.3 The Augmented Model

The augmented model is a trade-off between the above two extreme cases where both IP/MPLS and optical layers have separate control planes and routing instances, but allowing the exchange of some network information between the layers, such as reachability and/or summary of link state information (e.g. residual capacity), depending on a

necessary and specific agreement between the two layers.

2.2.4.4 Comparison of the Three Network Models

A summary of the main differences among the three models is shown in Table 2.1 [36]. All the three models have their advantages and disadvantages. In a peer model, common control and management plane allows avoiding duplication of the functions performed on these planes in each resource layer. Also, due to integration of all resources into a single transport plane, there is no need for standardization of the UNI interface between IP/MPLS routers and OXCs. On the other hand, integration of different client networks into a single transport plane is difficult. A single control plane in the peer model makes all the information freely accessible in the client domain, while in the overlay and augmented models, the optical network topology and the resource information are kept secure. They are more suitable for the case in which each layer is owned by different entity. Unlike the overlay model, the peer model supports dynamic routing that can either use only the existing lightpaths or open one or more lightpaths if found useful [37].

	Overlay	Augmented	Peer
Routing	Separated	Separated	Integrated
Network Information Exchanged	No information	Part or Summary	Full information
Signaling and control plane	Separated	Separated	Unified

Table 2.1: Network Models in IP/MPLS over WDM Networks

2.3 Survivability in Mesh Networks

As mentioned before, many mission-critical applications take place over the Internet, which requires high availability, reliability and QoS guarantees from the network. However, communication networks are subject to a variety of failures caused by natural disasters, wear out, and human errors, etc. Hence, being able to provision survivable services with guaranteed availability and QoS in real-time is also a key feature of the next generation networks [38].

Survivability is the capability of a network to maintain service continuity in the presence of faults within the network [39]. A critical component of network survivability technique is spare capacity allocation (SCA) problem in mesh network. The SCA problem is to decide how much spare capacity should be reserved on each link and where to pre-plan backup paths to protect traffic from a set of failures.

A number of survivability schemes have been proposed and extensively investigated in the past. In this section, we present a review of different survivability mechanisms in the IP/MPLS over WDM networks.

2.3.1 Survivability Schemes

In GMPLS-based mesh networks, there are two main survivability schemes: *protection* and *restoration* [40]. The major difference between the two is that in protection, a backup path or backup path segment is determined along which spare capacity is allocated at the time of connection setup or network design (i.e., prior to the failure), whereas in

restoration, it is dynamically determined along which spare capacity is allocated when the failure happens.

2.3.1.1 Protection

In mesh networks, protection schemes can be classified into link (local) protection, segment protection, and path protection. In *link protection*, the traffic is rerouted only around the failed link/node. In *segment protection*, a working path is divided into a sequence of segments, each segment is protected by a backup segment [6]. In *path protection*, the traffic is rerouted through a backup path once a link failure occurs on its working path. Path protection can be further classified into either failure-dependent or failure-independent. With failure independent path protection, the working path and backup path for a connection must be SRG-disjoint so that no single failure can affect both of these paths. With failure-dependent path protection, multiple backup paths that are not necessarily link/node-disjoint with a given working path are selected and which one to use depends on which link on the working path has failed, but rerouted traffic always goes through the source node (either the rerouting takes place at the source node or rerouted traffic loops back to the source node)

In a sense, link protection is similar to failure-dependent path protection in that in link protection, which detour to take also depends on which link has failed, except that link protection uses local rerouting. Segment-based protection schemes are also similar to failure-dependent path protection, but rerouted traffic only needs to go through the

node that starts a backup segment which protects the failed segment, as opposed to the source as in path protection (or the immediate upstream node of the failed link as in link protection). In general, both link and path protection can thus be considered as a special case of segment-based protection.

While path protection leads to efficient utilization of spare capacity, link protection provides fast recovery time. Segment protection can achieve high scalability and fast recovery time with a slight degradation in resource utilization.

Link, segment and path protection can be classified into dedicated or shared according to how the spare capacity is allocated for protection purposes. In *dedicated protection*, there are no sharing between backup resources, while in *shared protection*, spare capacity can be shared as long as their working paths are mutually diverse. The routers (OXCs) cannot be configured until the failure occurs if shared protection is used. Hence, recovery time in shared protection is longer but its resource utilization is better than that in dedicated protection.

2.3.1.2 Restoration

Restoration can also be classified into link, path and segment restoration depending on the type of rerouting [41]. In *link restoration*, when a failure occurs, the end nodes of the failed link dynamically find an alternative path for each connection that traverses the failed link. In *path restoration*, when a failure occurs, the source node of each connection that traverses the failed link is informed about the failure and find an alternative path

on an end-to-end basis. In *segment restoration*, when a failure occurs, the upstream node of the failed link detects the failure and find an alternative path to the corresponding destination node of each disrupted connection. The work in [41] has compared the performance tradeoff of these different restoration mechanisms under a distributed control and signaling system using GMPLS. It is shown that link restoration is most effective in terms of restoration time while path restoration is slowest.

2.3.1.3 Protection vs. Restoration

Figure 2.3 summarizes the classification of protection and restoration schemes. Generally, dynamic restoration schemes are more efficient in utilizing the network capacity because they do not allocate spare capacity in advance, and they provide resilience against one or multiple failures as long as the destination is still reachable. However, they cannot guarantee the recovery time, and/or the amount of information loss for real-time applications, making them unsuitable for mission-critical applications.

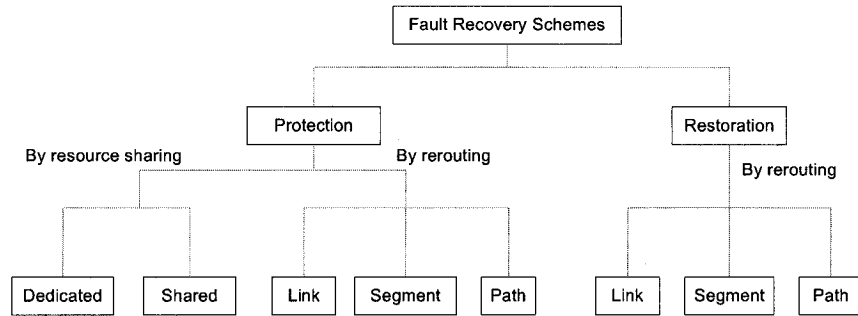


Figure 2.3: Protection and restoration schemes in mesh networks

On the other hand, protection schemes are designed to cope with a predefined number of simultaneous failures. They are less bandwidth efficient than restoration schemes, but have faster recovery time and can guarantee the quality of recovery from disrupted services. Therefore, protection mechanisms are the best techniques to improve the availability of the supported services in GMPLS networks.

2.3.1.4 Shared Backup Path Protection

As mentioned before, path-oriented survivability schemes are more capacity efficient than their link-based and segment-based counterparts. There are two types of path protection schemes: shared backup path protection (SBPP) and dedicated backup path protection (DBPP). SBPP is very similar to DBPP in that working traffic between an S-D pair can be recovered over a predefined failure-disjoint path. But in SBPP, spare capacity on the backup paths can be shared by backup paths whose working paths do not share a common failure with each other, greatly reducing capacity redundancy.

SBPP can be implemented using the MPLS/GMPLS protocol suite [42]. The process is as follows. Each node maintains a local network status database, which records the complete information of network resource usage and spare capacity sharing relationships on each link. Database synchronization relies on the Link State Advertisement (LSA) but extended to contain TE resource data. Whenever there is any network status change, all the nodes in the network are notified via an LSA flooding message, and update their local databases of global network state. On the arrival of a connection request, the

source node checks the local copy of the network state database to see if the network can currently support a new working and backup path pair to the requested destination. If so, the source node uses RSVP-TE [27] or CR-LDP [28] signaling to establish the working and backup paths. LSA messages are then disseminated network-wide to update the network state and spare capacity sharing databases at other nodes. The connection is released using the same signaling protocols. In case of a failure, the source node detects the failures and uses the same signaling protocols to request cross-connections along the backup path to activate the backup path in real time.

In summary, SBPP has several desirable features :

- “It allows protection to be arranged (or not) at the discretion of the user and lets the user know the route of their backup path in advance if failure occurs”. [43]
- End-to-end rerouting gives customers control on activating the backup paths for their affected services.
- It achieves a very efficient protection-to-working capacity ratio.

Therefore, in this thesis, we are particularly interested in the approach of allocating spare capacity using SBPP. Though network redundancy is reduced to some degree by sharing spare capacity, in this thesis we are interested to have redundancy minimization as an optimization criterion with the consideration of connections’ availability constraints.

2.3.2 Single-layer and Multilayer Survivability

All the recovery mechanisms presented in Section 2.3.1 are available in both the IP/MPLS layer and the optical layer. They could be employed only at a single layer or at both layers.

2.3.2.1 Recovery in the Optical Layer Only

In this case, all the recovery actions are performed in the optical layer. Because of the coarser switching granularity of optical layer, this recovery approach is simpler in the number of affected paths to reroute, and failures do not propagate to the upper layer. This will lead to reduced signaling overhead to notify the end nodes of the failed lightpaths and activate backup lightpaths and guarantee fast recovery within a few tens of milliseconds [37].

However, optical layer recovery has some limitations and drawbacks [44]. Firstly, the total capacity investment for restorability can be expensive because of the coarser restoration granularity. Secondly, this recovery strategy cannot handle problems that occur due to failures in the upper layer, e.g., nodes failures in the IP/MPLS layer can only be recovered by the actions of peer-level network elements. In case of a node failure in the optical layer, the upper layer node might be isolated.

2.3.2.2 Recovery in the IP/MPLS Layer Only

In this case, all the recovery actions are performed in the IP/MPLS layer. Failures in both layers (IP/MPLS or optical layer) can be handled by survivability schemes in the IP/MPLS layer. Another advantage is, finer granularity of traffic in the IP/MPLS layer allows for the implementation of a resilience differentiated approach which protects different traffic flows with different recovery granularity, and QoS granularity [45]. For example, the resilience can be employed only for individual services requiring a high availability. The services with lower availability requirements could be unprotected or partially protected. Such differentiated resilience results in a more cost effective network design and traffic engineering in comparison to the previous scenario. Therefore, it is reasonable for an Internet service provider (ISP) to provide the required network survivability using only resilience mechanisms in the IP/MPLS layer.

A drawback of this approach is that in case of a failure in the optical layer, a large number of individual flows in the IP layer will be disrupted and many recovery actions may be needed [46].

2.3.2.3 Layer Interworking

Both of the above two single-layer recovery approaches have their pros and cons. The advantages of these approaches can be combined by allowing recovery mechanisms of different layers to cooperate in recovering from failures. However, the presence of resilience mechanisms in multiple layers leads to a contention. This contention may result in sub-

optimal recovery and network resources could be inefficiently used. Hence, the multilayer resilience requires the interaction of recovery mechanisms present in the different layers. However, the interworking between layers requires some rules in order to ensure efficient recovery process.

Three escalation strategies that define how the layers and the recovery mechanisms within those layers react to different failure scenarios are presented in [46]: uncoordinated, sequential, and integrated escalation.

In *uncoordinated approach*, the recovery schemes are deployed in the multiple layers respectively with any coordination at all. This results in parallel recovery actions at different layers. The advantage of this approach is that it is simple from an implementation and operational point of view. The drawback is that in case of a failure in the optical layer, both recovery mechanisms occupy spare resources during the failure, although one recovery scheme occupying spare resources would have been sufficient. This situation could even be worse with recovery mechanisms in different layers locking each other in some cases or resulting in routing instabilities.

In comparison with the uncoordinated approach, *sequential approach* is a more efficient escalation strategy. Here, if the current network layer cannot recover the affected traffic within the predefined time, the responsibility for recovery is handed over to the next layer. There are two strategies to coordinate the recovery mechanisms of different layers. In *bottom-up escalation*, the recovery starts in the lowest detecting layer and escalates upwards. The advantage is that the recovery is performed at the appropriate

granularity. In *top-down escalation*, recovery actions are first initiated in the highest possible (IP) layer. Only if the higher layer cannot restore all affected traffic, the lower layer recovery mechanisms are triggered. An advantage of this approach is that a higher layer can differentiate traffic with respect to service types. The drawback is that it is not easy for a lower layer to detect if a higher layer can restore all the affected traffic, and the implementation is very complex.

In *Integrated approach*, a common integrated recovery mechanism is employed across all layers. This is the most flexible, but also most complicated recovery approach among the three escalation strategies.

2.4 Connection Availability Analysis in Mesh Networks

2.4.1 Mathematical Definitions

2.4.1.1 Reliability

The reliability of a system is defined as the probability that the system will perform its intended function during a defined period [47]. The reliability function can be defined as a function of the time T during which no system failures happen:

$$R(T) = P\{\text{no failure in } [0, T]\} \quad (2.1)$$

which can also be expressed in terms of the *failure density function* $f(t)$ as follows:

$$R(T) = 1 - \int_0^T f(t)dt \quad (2.2)$$

The function $f(t)$ is in fact the probability density function of *the time to failure* random variable. It is an instantaneous rate of failure, therefore integrating $f(t)$ over a certain period gives the probability that the first failure will occur in that time period. Conversely, by differentiating Equation (2.2) one can express $f(t)$ in terms of $R(T)$:

$$f(t) = \frac{d}{dt}R(t) \quad (2.3)$$

The expectation of *the time to failure* gives the *mean time to failure* (MTTF), which is a useful measure in availability analysis in the following section.

$$MTTF = E(\text{the time to failure}) = \int_0^\infty (t \cdot f(t))dt \quad (2.4)$$

2.4.1.2 Availability

The concept of availability is related to repairable systems. It is defined as “the probability of the system being found in the operating state at some time t in the future given that the system started in the operating state at time $t = 0$ and given that failures and down-states occur but maintenance or repair actions always return the system to an operating state.” [47]

The availability of a system is a function of time that starts from 1 and usually stays at a high level shortly after the system starts operating and then decreases to eventually reach a steady state in which repairs compensate for failures and maintain the availability at a certain constant level. Therefore, the steady state availability A of a system can be expressed as the fraction of time the system is up over a long period of time T [48]:

$$A = \lim_{T \rightarrow \infty} \left\{ \frac{\text{up time}}{T} \right\} \quad (2.5)$$

Based on Equation (2.5), [48] gives the derivation of a very useful expression of A :

$$A = \frac{MTTF}{MTTF + MTTR} \quad (2.6)$$

Equation (2.6) is sometimes replaced by the following expression:

$$A \approx \frac{MTBF}{MTBF + MTTR} \quad (2.7)$$

where $MTBF$ is the mean time between failures and $MTTR$ is the mean time to repair. *Time between failures* in $MTBF$ refers to the time between the occurrence of failures, whereas *time to failure* in $MTTF$ is the time between the repair of a failure and the occurrence of the next failure as shown in Figure 2.4. Therefore,

$$MTBF = MTTF + MTTR \quad (2.8)$$

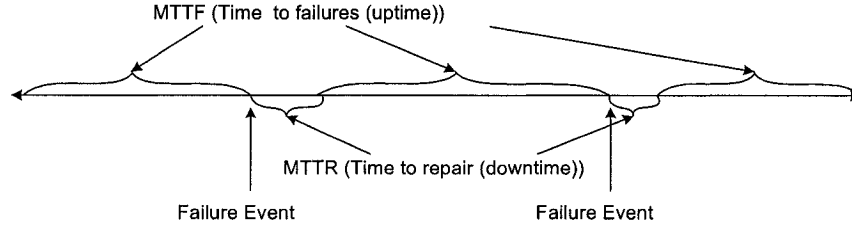


Figure 2.4: An illustration of MTTF and MTTR

Strictly, Equation (2.6) gives the accurate result. However, with typical values of $MTTF$, $MTBF$ and $MTTR$, Equations (2.6) and (2.7) give almost the same results.

Sometimes, instead of working on availability values, it is easier to work with unavailability values of a system which can be expressed as the complement of the availability:

$$U = 1 - A \quad (2.9)$$

It can also be expressed in terms of $MTTF$, $MTBF$ and $MTTR$:

$$U = \frac{MTTR}{MTBF} \approx \frac{MTTR}{MTTF} \quad (2.10)$$

2.4.2 Availability Analysis

The mathematical definitions we introduced above give the instruments to calculate the reliability and availability parameters of a complex *system*. A general method is to develop a logical block diagram with the different *functional blocks* that compose the

system. The block diagram can often be reduced by applying some simplifications for elements in series and elements in parallel.

In our study, the system to be characterized is a network component, a path or a connection in a mesh network. Our availability analysis is based on the following typical assumptions [10]:

1. A network component is either available or unavailable.
2. Different network components are mutually failure-independent and failures randomly occur in time, independently from the components age.
3. For any network component, the variables $MTTF$ and $MTTR$ are independent memoryless processes with known mean values.

2.4.2.1 Network Component Availability

According to previous assumptions and introduced availability theory, the availability of a network component (e.g., a router or a link) which is denoted as A_i can be expressed as follows [48]:

$$A_i = \frac{MTTF}{MTTF + MTTR} \quad (2.11)$$

2.4.2.2 Availability of an End-to-End Path

A path c can be represented by a *series* of N network components on which it is routed. The connection is available only when all the components along its route are available.

Let A_i denote the availability of i th component, then the end-to-end path availability which is denoted as A_c can be computed as follows:

$$A_c = \prod_{i=1}^N A_i \quad (2.12)$$

For realistic systems, $A_i \approx 1$ and $U_i \ll 1$, where U_i denotes the unavailability of the i th component. A useful approximate equation is usually employed to evaluate the connection unavailability U_c [13]:

$$\begin{aligned} U_c &= 1 - A_c \\ &= 1 - \prod_{i=1}^N (1 - U_i) \\ &= \sum_{i=1}^N U_i - \sum_{i,j=1, i \neq j}^N U_i \cdot U_j + \sum_{\substack{i,j,k=1 \\ i \neq j, i \neq k, j \neq k}}^N U_i \cdot U_j \cdot U_k - \dots \\ &\approx \sum_{i=1}^N U_i \end{aligned} \quad (2.13)$$

2.4.2.3 Availability of Path-Protected Connections

It is well known that a protection scheme helps to improve a connection availability since traffic on the failed working path (link/segment) is quickly switched to the backup path (segment). For example, in a network which is designed to be 100% restored upon any single failure, a connection equipped with a failure-disjoint backup path has 100% availability in the presence of any single failure. Nevertheless, when multiple failures is considered, the connection availability depends intimately on the locations of the

failures, how much backup resources are reserved (i.e, single backup path or multiple backup paths, full restoration or partial restoration), and how the backup resources are allocated (i.e, dedicated or shared).

As described in Section 2.3.1.4, DBPP and SBPP are two types of path protection schemes. In Figure 2.5, a DBPP connection c is carried by one working path w and protected by one protection path p which is node-disjoint with w . Path w consists of N components and path p consists of M components. In DBPP, the state of availability of the two paths does not depend upon the state of any other path, that is, each DBPP connection can be treated separately from all the other connections in the network. The connection c is up when either the working, the protection or both the paths are available. The connection c is down only when both w and p are unavailable. As we said earlier, if only a single failure occurs in the network, the end-to-end availability of connection c will be 100%. If considering multiple failures could happen, A_c can be exactly computed as follows:

$$A_c = 1 - (1 - A_w) \cdot (1 - A_p) \quad (2.14)$$

where A_w and A_p denote the availabilities of w and p , respectively, and can be computed by Equation (2.12). If we work with unavailability values of the network components, the total unavailability of connection c can be obtained by considering the *parallelism* of

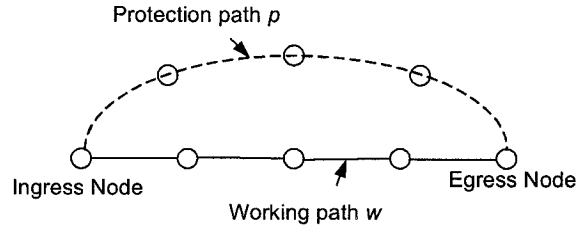


Figure 2.5: An example of dedicated path protection

working and protection paths:

$$U_c = U_w \cdot U_p = \sum_{i=1}^N U_i \cdot \sum_{j=1}^M U_j \quad (2.15)$$

A connection may employ multiple backup paths to increase its availability. If all backup paths are disjoint and dedicated to this connection, the connection availability can be derived following the similar principles in Equations (2.14) and (2.15).

In a typical implementation of the SBPP scheme, a connection is carried by a single working path and protection path as in the DBPP case, but the spare capacity along the protection path can be shared by other protection paths whose working paths do not share a common failure. Such a connection will have 100% availability in the presence of any single failure. However, in the case of multiple failure scenarios, the analysis of connection availability gets more complicated.

For example, in the network shown in Figure 2.6, W_1 , W_2 and W_3 are three working paths between three S-D pairs AC , CG , and FG respectively. We assume each connection

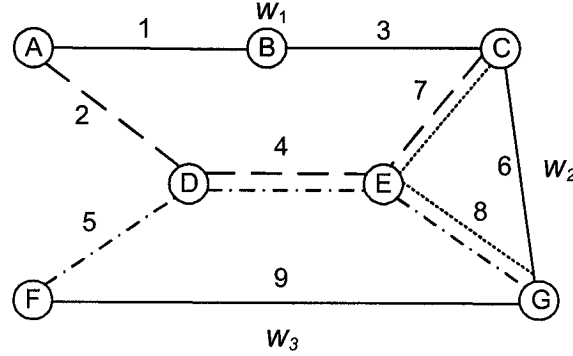


Figure 2.6: An example of shared backup path protection

takes one unit bandwidth. W_1 and W_2 share spare capacity on link 7; W_1 and W_3 share spare capacity on link 4; and W_2 and W_3 share spare capacity on link 8. Therefore, only one unit spare capacity needs to be reserved on link 4, 7 and 8 respectively. In the case that only link 1 fails, the traffic on W_1 will be switched onto its backup path which traverses through links 2, 4, and 7. If link 9 fails first, and link 1 fails before link 9 is repaired, the connection between S-D pair FG will get the spare capacity on link 4 while the connection between S-D pair AC will be disrupted. However, if two units bandwidth are reserved as spare capacity on link 4, both connections could survive.

Obviously, when multiple failures are considered, the availability impairment in SBPP is worse than that in DBPP due to resource sharing along the backup paths. Whether or not the spare capacity is available to the protection path of a specific working path depends not only on the physical availability of the links taken by its protection path, but also on the availability of the other working paths with their protection paths sharing

the spare capacity along these links, the location of the failures, and how much backup resources are reserved. Therefore, it is very difficult to give a closed-form equation for computing the availability of an SBPP connection. In order to assess the availability of the connection with sufficient accuracy, combinations of multiple failure events should be enumerated to take into account the availability impairment of each failure situation. The list of such combinations becomes rapidly lengthy as the shared spare capacities and the number of sharing connections increase. An efficient availability model for SBPP connections is needed and will be proposed in the next chapter.

2.5 Summary

A network architecture of IP/MPLS over WDM is expected to form the base of next generation internet backbone networks. This chapter has first provided an overview of the control planes, network models, and survivability schemes in IP/MPLS over WDM networks. Specifically, we have reviewed the important characteristics of MPLS and GMPLS based control planes. The pros and cons of the different approaches for providing recovery in multilayer networks have been examined. The important definitions and concept of service (connection) availability analysis have been introduced. We have also provided the availability analysis for connections with or without path protection.

Chapter 3

Availability-Aware Spare Capacity Allocation and Reconfiguration

3.1 Introduction

In this chapter, we first provide a review of the literature in the area of availability evaluation and modeling for connections with SBPP and partial protection. Then, the system and the problem formulations are presented. We then propose two availability-aware spare capacity allocation (SCA) models based on two policy-based availability models for SBPP connections. A novel spare capacity reconfiguration (SCR) architecture is then designed based on the proposed SCA models.

3.2 Related Work

To the best of our knowledge, the majority of previous work in the area of availability evaluation and modeling for connections with SBPP has focused on the design in the optical layer. In [1], the network restorability has been studied in a network designed for single failure when double link failure occurs in mesh networks with link-protection. The study in [10] has conducted availability analysis for a connection with no protection, dedicated protection and SBPP considering all multiple failure combinations in the network, and has introduced an ILP-based availability-aware connection provisioning under the static traffic. In [11] and [12], the authors have proposed an availability-aware provisioning algorithm for dynamic traffic to address the availability constraint under multiple failures by manipulating a routing metric. The study in [49] has evaluated the connection availability of a number of protection techniques, including shared protection for which they have adopted approximation by considering no more than two simultaneous failures. The authors in [13] have proposed a heuristic algorithm for designing optical networks with a set of protected static connections to maximize the availability of each connection and minimize the deployment cost. The study did not consider the availability impairment due to spare capacity sharing for connections with SBPP.

A suite of availability analysis and capacity design methods have been introduced in [50], where dual-failure events in optical networks supporting SBPP connections are considered. The authors in [50] have proposed to improve the E2E availability of SBPP connections by limiting the number of service paths sharing the same unit of spare

capacity. In [51], two traffic grooming algorithms have been introduced to provision dynamic bandwidth with guaranteed availability, where dedicated or shared protection can be used, and all possible multiple failures are considered. The study in [52] have collected availability parameters for various network components and compared the E2E availability of connections under different resilience mechanisms. In [14], a matrix-based approach has been introduced in estimating the unavailability of dual-failure patterns using a Markov chain model, where the sequence of failures in the failure patterns is considered. The authors in [15] have proposed a dynamic SBPP connection provisioning algorithm in optical networks based on the matrix-based approach defined in [14]. The studies in [14] and [15] have investigated the case that either 100% or none of the working capacity of a connection is restored in presence of a failure event.

Among the studies dealing with IP/MPLS layer protection and restoration, [53] has introduced an off-line strategy of searching for a primary path with improved availability for static traffic, where all possible multiple failures are considered. The study in [54] has exploited a strategy that integrates the knowledge of protection in the optical layer to dynamically provision E2E availability-guaranteed services in the IP over optical networks. It is noticeable that only blocking due to insufficient service availability is considered but due to insufficient bandwidth. The study in [55] has introduced a method to calculate the availability of a dedicated backup path protected LSP which considered the availability of link buffer in the MPLS layer. The calculation of E2E connection availability has been introduced in [56] based on a number of thumb rules, such as the

E2E availability under a protection scheme for all single failures or for all single and dual failure events.

In all the abovementioned studies, the source node of a connection switches 100% of its bandwidth over to the protection path when failure occurs to the working path. This policy is necessary for the restoration of connections with indivisible bandwidth such as in the optical layer. In the IP/MPLS layer, an LSP may support numerous independent service sessions such that dropping any/some of them would not affect the others. Thus, it could be unnecessary to require the interrupted working LSP to be either 100% restored or non-restorable.

The concept of partial restoration has been investigated in [57][58][59]. The study in [57] has introduced partial protection using the deterministic Quality of Protection (QoP) paradigm. The study in [58] has conducted extensive simulations and concluded that the partial restorability could lead to smaller resource consumption than that in the full restorability case. The authors in [59] have demonstrated that partial restorability on the video streams in SONET/SDH rings leads to smaller capacity demand. The research has concluded that the consumed resources are a linear function of the fraction of restorability. To our best knowledge, all the studies on partial restoration have never touched the availability evaluation, and no information has been provided on how the source node randomly drops part of the working bandwidth in the restoration phase.

It is envisioned that making the working LSP partially restorable would greatly improve the design flexibility, restoration granularity, and capacity efficiency in the event

that the E2E availability constraint on each connection is the ultimate goal of network operation instead of whether the LSP can be 100% restored in the presence of some number of simultaneous failures.

3.3 System Formulation

In this work, we concentrate on a survivable IP/MPLS network. In our model, we assume the only information available to the IP/MPLS layer regarding the underlying optical network are the SRGs. We only consider the availability impairment upon a connection in the IP/MPLS layer due to any physical component failure. No service unavailability caused by traffic congestion in the IP/MPLS layer is considered. We also assume that only the IP/MPLS layer provides protection for each working LSP for availability enhancement.

Without loss of generality, only failures on each IP link are considered, while each IP router is taken as perfect. This assumption is reasonable since in general heavy redundancy and extremely short recovery time can be achieved for an IP router which is usually located in a city. Under this assumption, an SRG is defined as a set of IP links that are simultaneously struck by a common failure event. We assume each SRG is either available or unavailable due to a failure, which happens independently. The repair time and the time to failure of each SRG are memoryless, exponentially distributed random processes with constant Mean-Time-to-Failure (MTTF) and Mean-Time-to-Repair (MTTR). Each IP link can be associated with an SRG identifier that

allows the IP/MPLS layer to find out which IP links are routed disjointedly in the optical layer.

According to the analysis in Section 2.4, the E2E unavailability of a connection with SBPP is evaluated directly by enumerating the failure patterns that affect the connection, where the stationary probabilities of all considered failure patterns are summed up. The availability of a connection is the complement of its unavailability.

Let R denote a specific set of failure patterns in the network which are considered in evaluating the availability of a connection. Let the availability of connection c regarding R be denoted as A_c^R , and the stationary probability of failure pattern $r \in R$ be denoted as π_r . We also define a set R' that contains only the failure patterns that make connection c not restorable at all, i.e,

$$R' = \{r | c \text{ is 100\% not restorable, } r \in R\}$$

Obviously, $R' \subseteq R$. Therefore, A_c^R can be evaluated as

$$A_c^R = 1 - \sum_{r \in R'} \pi_r$$

It is obvious that a protection scheme can achieve *full restorability* (or $A_c^R = 1$) in case $R' = \emptyset$. An example for a protection scheme with full restorability is that a working path is protected by a single SRG-disjoint backup path under the single failure scenario. In this case, R contains all the failure patterns with a single SRG, and the derived E2E

availability of the connection regarding such R is 100%.

A protection scheme with *partial restorability* is in contrast with the case of full restorability such that $R' \neq \emptyset$ and $A_c^R < 1$. By taking the previous example (that a working path is protected by a single SRG-disjoint backup path), when R contains not only the failure patterns with a single SRG but also the ones with multiple SRGs, the E2E availability of such a connection would become less than 100%. From the service point of view, the working LSP is *partially restorable* under a specific protection scheme if the expected restorable bandwidth in presence of $\forall r \in R$ is a proportion of the bandwidth provisioned by the working LSP.

Partial restorability of a working path can be achieved in the following two scenarios:

1. The working path is partially protected by one or multiple backup paths and/or path segments such that only a specific set of failure patterns defined in R is restorable. In this case, equipping each backup path/path segment can remove the corresponding failure patterns from R' . It is clear that by adding more effective backup paths/path segments to the working path, less failure patterns will be contained in R' , and A_c^R will be further increased. For example, in a typical shared backup path protection scheme, if R contains single, dual and triple failure patterns, the set of non-restorable failure patterns R' for a connection includes dual and triple failure patterns. However, if two backup paths are provisioned for each working path, dual failure patterns can be removed from R' and the connection availability will increase.

2. The spare capacity allocated along a backup path/path segment is a fraction of the total working bandwidth. In this case, the switching node of the backup path/path segment is equipped with the intelligence that switches and restores a fraction of the total working bandwidth while the rest of the bandwidth is disregarded.

Partial restorability in the second scenario is feasible for a working LSP provided with the following two conditions:

- The bandwidth of the LSP is divisible in the restoration plane;
- An LSP supports numerous independent service sessions, and dropping some of them would not affect the integrity of the others.

For the first condition, since the MPLS control plane supports finer switching granularities, thus, the restoration mechanisms in the IP/MPLS layer can provide finer recovery granularities and provide efficient and flexible resource usage by restoring a fraction of the bandwidth of a working LSP. The second condition is a common case in the MPLS networks since an LSP is potentially aggregated with a number of independent LSPs or service sessions with the same FEC. Under such a circumstance, it is relevant to assume the required intelligence at the switching nodes of an LSP that perform partial restoration in the second scenario.

With the two scenarios for partial restoration, the spare capacity allocation problem can be engineered with better flexibility and more design granularities in the effort of meeting the E2E availability requirement of each connection request. Since this work is

interested in optimizing the spare capacity consumption for connections equipped with a single shared SRG-disjoint backup path, we focus on the second scenario of partial protection.

3.4 Problem Formulation

Let the network topology be represented by $G(V, E)$, where V is the set of nodes, E is the set of IP links. Since it is almost impossible (and also unnecessary) to consider and enumerate all possible failure patterns in the network when evaluating the E2E availability of SBPP connection, the set of failure patterns considered in the availability evaluation could be made to include the failure patterns only with relatively large stationary probabilities while the ones with fairly low probabilities are disregarded. In this work, we assume that no more than two simultaneous SRG failures will happen since the probability of M simultaneous failures drops dramatically when $M > 2$ [60].

For simplicity, we assume each SRG takes only one IP link. Therefore, R contains all single and dual failure patterns in the network, where $r_m \in R$ denotes the single failure pattern that only link m fails and all other links operate in the network, and $r_{mn} \in R$ denotes the simultaneous failures on link m and n with link m failed first followed by the failure of link n . In this case, the number of all failure patterns (denoted as $|R|$) in the network is $|E|^2$. The stationary probabilities of failure patterns r_m and r_{mn} are denoted as π_m and π_{mn} , respectively. Given the values of MTTF and MTTR of each link in the network, the probabilities of all the failure patterns can be derived by solving the

Markov chain introduced in [14]. Please see Appendix A for this computation.

Let the network be launched with a set of connections denoted as C . All connections are provisioned with a working and link-disjoint shared protection paths. Each connection request $c \in C$ is associated with a tuple $\langle s_c, d_c, A_{c,SLA}, B_c \rangle$, where s_c is the source, d_c is the destination, $A_{c,SLA}$ is the availability requirement of connection c specified in the SLA, and B_c is the bandwidth requirement. Let w_c and p_c denote the working and protection path of connection c , respectively, and let W_c and P_c denote the set of links along w_c and p_c , respectively.

For a generic connection c , all the failure patterns in R can be divided into four subsets as follows:

- $R_{\overline{wp}}^c$: the set of failure patterns which interrupt both working and protection paths of c , i.e., $R_{\overline{wp}}^c = \{r_{mn} | (m \in W_c, n \in P_c) \vee (m \in P_c, n \in W_c)\}$
- $R_{\overline{w_1p}}^c$: the set of failure patterns whose first failure disrupts the working path of c and second failure does not affect the protection path of c , i.e., $R_{\overline{w_1p}}^c = \{(r_m | m \in W_c) \cup (r_{mn} | m \in W_c, n \notin P_c)\}$
- $R_{\overline{w_2p}}^c$: the set of failure patterns which interrupt the working path of c by the second failure while the first failure does not affect the protection path of c , i.e., $R_{\overline{w_2p}}^c = \{r_{mn} | m \notin (W_c \cup P_c), n \in W_c\}$
- R_w^c : the set of failure patterns containing all the failure patterns that do not disrupt the working path of c , i.e., $R_w^c = \{r | r \in (R - R_{\overline{wp}}^c - R_{\overline{w_1p}}^c - R_{\overline{w_2p}}^c)\}$

To better optimize the backup resource allocation, partial protection is adopted. We define *protection level* of a connection as the percentage of the working bandwidth to be restorable by the protection path of this connection once the working path is interrupted. In this work, the protection level of a connection can be *failure-independent* (FID) or *failure-dependent* (FD). The FID policy takes the same protection level for the working path regardless of the location of the failure pattern. On the other hand, the FD policy defines a specific protection level along the given protection path corresponding to each failure pattern that affects the working path.

3.5 Proposed Spare Capacity Allocation Models

In this section, the proposed models for availability-aware spare capacity allocation (SCA) are introduced, where two novel linear programs (LPs) are formulated according to the adopted policy in routing and spare capacity sharing. The availability model of each connection is taken as a constraint in the SCA formulations.

Let $S_{j,r}$ denote the total amount of spare capacity required to be allocated on link j when failure pattern r occurs in order to guarantee the required protection level of each connection. The failure patterns considered in the study could be either a single-failure event r_m or a dual-failure event r_{mn} , where m and n are two SRGs. Let V_j denote the spare capacity allocated on link j . Given the knowledge of all the working paths W_c and P_c for all $c \in C$, we can formulate the spare capacity allocation problem into an LP in the scenario of Failure-Independent (FID) and Failure-Dependent (FD), which are

presented in the following sections.

3.5.1 Failure-Independent (FID) SCA

3.5.1.1 Mathematical Formulation

In the FID policy, the protection level of connection $c \in C$ is represented by θ_c regardless of the failure pattern. The objective is to minimize the total spare capacity that needs to be reserved for all the active connections such that their availability requirement can be met:

$$\min \sum_{j \in E} V_j \quad (3.1)$$

V_j is targeted to guarantee only the spare capacity required by the connections interrupted by any single failure pattern with the corresponding protection level for each connection is subject to the availability requirement of the connection, which can be expressed by the following constraint:

$$V_j = \max_{r_m} S_{j,r_m} \quad \forall j \in E, r_m \in R \quad (3.2)$$

where S_{j,r_m} satisfies the following constraint:

$$S_{j,r_m} = \sum_{\substack{\forall c \text{ s.t. } m \in W_c, j \in P_c}} B_c \cdot \theta_c \quad \forall j \in E, r_m \in R, m \neq j \quad (3.3)$$

The function $\max$ in Equation (3.2) asserts that for link $j \in E$, V_j is always greater than or equal to the maximum spare capacity required on link j by any single failure pattern r_m .

With constraints (3.2) and (3.3) defined, the failure patterns belonging to $R_{\overline{w}p}^c, R_{\overline{w}1p}^c$ and $R_{\overline{w}2p}^c$ may have different availability impairment to connection c instead of simply blocking the service.

For the failure patterns belonging to $R_{\overline{w}p}^c$, both w_c and p_c are unavailable, thus, the availability impairment is 100%. For the failure patterns belonging to $R_{\overline{w}1p}^c$, w_c is interrupted by the first failure, and gets its reserved spare capacity $B_c \cdot \theta_c$ on its protection path. Thus, the availability impairment is $1 - \theta_c$.

When a dual-failure pattern r_{mn} belonging to $R_{\overline{w}2p}^c$ occurs, the situation becomes more complicated. Let C_{j,r_n} denote the set of all the connections whose working paths are disrupted by r_n and protection paths traverse through link j . The spare capacity required on link j by all the working paths interrupted by r_{mn} is denoted as $S_{j,r_{mn}}$ which satisfies the following equation:

$$S_{j,r_{mn}} = S_{j,r_m} + S_{j,r_n} - \sum_{\forall c \text{ s.t. } m \in W_c, j \in P_c} B_c \cdot \theta_c \quad \forall r_{mn} \in R, j \neq m, j \neq n \quad (3.4)$$

In this case, the following two scenarios introduce different availability impairments upon connection c :

1. If $V_j \geq S_{j,r_{mn}}$ on all links $j \in P_c$, all connections in C_{j,r_n} get their required spare

capacity, and there is no contention upon the spare capacity between them. The availability impairment on c is $1 - \theta_c$.

2. If $S_{j,r_m} \leq V_j \leq S_{j,r_{mn}}$ on all links $j \in P_c$, the residual spare capacity after recovering the connections affected by the first failure r_m will not be enough for the required protection level for the connections in C_{j,r_n} . Due to the insufficiency, the source nodes of these connections drop their traffic proportionally according to their protection levels and the available spare capacity. The lower and upper bound on the availability impairment for c would be $(1 - \theta_c, 1]$.

Thus, when dual-failure pattern r_{mn} belonging to R_{w2p}^c occurs, the availability impairment on connection c will be in a range of $[1 - \theta_c, 1]$. To simplify the availability model, we assume that the availability impairment on c is always $1 - \theta_c$ no matter which scenario happens. Our approximation yields an upper bound on the E2E availability of connection c , however, has the advantage that the spare capacity contention does not have to be included into the model and the SCA formulation is linear. On the other hand, the derived protection level for each connection from FID-SCA model is a lower bound on the actual protection level that could satisfy the availability requirement of the connection.

With the above approximation, the E2E availability of connection c denoted as A_c

can be calculated as follows:

$$A_c = 1 - \sum_{r \in R_{wp}^c} \pi_r - \sum_{r \in (R_{w1p}^c \cup R_{w2p}^c)} (1 - \theta_c) \cdot \pi_r \quad \forall c \in C \quad (3.5)$$

Obviously, in order to meet the availability requirement of connection c , the following constraint is required:

$$A_c \geq A_{c,SLA} \quad \forall c \in C \quad (3.6)$$

where θ_c is under the constraint:

$$0 \leq \theta_c \leq 1 \quad \forall c \in C \quad (3.7)$$

The objective function Equation (3.1) along with Equations (3.2), (3.3), (3.5), (3.6) and (3.7) constitute the formulation of FID-SCA.

3.5.1.2 Complexity

In FID-SCA formulation, the variables to be solved are V_j for each $j \in E$ and θ_c for each $c \in C$. Therefore, the number of variables grows as $O(|E| + |C|)$. Since the spare capacity on each link is reserved for all the connections disrupted by any single failure pattern and each connection must hold the availability constraint, the number of constraints grows as $O(|E| \times (|E| - 1) + |C|)$.

3.5.2 Failure-Dependent (FD) SCA

3.5.2.1 Mathematical Formulation

With the FD policy, we let $\theta_{c,r}$ denote the protection level of connection c in the occurrence of failure pattern r . The objective of FD-SCA is the same as that in FID-SCA, while the spare capacity sharing on link j is among all failure patterns in R . In other words, V_j is designed to serve as the minimal required spare capacity under any failure pattern for all the connections whose working paths are disrupted by the failure pattern and protection paths traversing through link j , i.e.,

$$V_j = \max S_{j,r} \quad \forall j \in E \quad (3.8)$$

where $S_{j,r}$ depends on the protection levels of those interrupted connections:

$$S_{j,r} = \sum_{\forall c \text{ s.t. } j \in P_c, r \in (R_{w_1p}^c \cup R_{w_2p}^c)} B_c \cdot \theta_{c,r} \quad \forall j \in E, r \in R \quad (3.9)$$

With constraints in Equations (3.8) and (3.9), the availability impairment due to the occurrence of failure patterns belonging to $R_{\overline{w}p}^c$ is 100%. In the occurrence of failure patterns belonging to $R_{w_1p}^c$ or $R_{w_2p}^c$, w_c is restored in a protection level $\theta_{c,r}$ since spare capacity sharing is among all failure patterns. Thus, the availability impairment is simply

$1 - \theta_{c,r}$, and the availability model of connection c is as follows:

$$A_c = 1 - \sum_{r \in R_{wp}^c} \pi_r - \sum_{r \in (R_{w1p}^c \cup R_{w2p}^c)} (1 - \theta_{c,r}) \cdot \pi_r \quad \forall c \in C \quad (3.10)$$

A_c is subject to the availability requirement specified in the SLA as Equation (3.6), where $\theta_{c,r}$ in Equations (3.9) and (3.10) is subject to the following constraint:

$$0 \leq \theta_{c,r} \leq 1 \quad \forall c \in C, r \in (R_{w1p}^c \cup R_{w2p}^c) \quad (3.11)$$

The objective function Equation (3.1) along with Equations (3.8), (3.9), (3.10), (3.11) and (3.6) constitute the formulation of FD-SCA.

3.5.2.2 Complexity

In FD-SCA formulation, the variables to be solved are V_j for each $j \in E$ and $\theta_{c,r}$ for each $c \in C$. The number of variable $\theta_{c,r}$ for connection c depends on the set of links traversed by the working path and protection path of the connection. A connection $c \in C$ is affected by $|W_c|$ single failure patterns and $2 \times |W_c| \times |P_c|$ dual-failure patterns. Therefore, the number of variables of FD-SCA formulation grows as $O(|E| + \sum_{c \in C} (|W_c| + 2 \times |W_c| \times |P_c|))$. Since the minimal required spare capacity is shared among all failure patterns in R , the number of constraints grows as $O(|E| \times (|E| - 1)^2 + |C|)$

Obviously, FD-SCA formulation has much more constraints and variables than FID-SCA formulation.

3.6 Proposed Spare Capacity Reconfiguration Architecture

Based on the proposed E2E availability models and SCA formulations in the previous section, an availability-aware dynamic provisioning strategy for SBPP connections is presented in this section, where *spare capacity reconfiguration* (SCR) is performed by solving the policy-based SCA formulations to reconfigure the spare capacity along each link.

In the proposed dynamic SCR framework, we adopt a centralized *network management system* (NMS) to conduct the network-wide reconfiguration. The NMS is assumed to have full knowledge of per-flow information (i.e., the existing working and protection path-pairs, along with their availability requirements). Each newly arrived connection request c is associated with a tuple $\langle s_c, d_c, A_{c,SLA}, B_c \rangle$ as described in Section 3.4.

A *network event* is defined as an event where the network traffic distribution differs from the last network event by an amount larger than a predefined threshold due to the traffic variation [61]. In this work, the threshold of traffic variation is simply defined as a specific number of connection setup or tear-down. A reconfiguration process is initiated right after a network event, and is called a “success” if the reconfiguration process is completed before the advent of the next network event. Thus, the necessary condition for a successful reconfiguration process is that the computation time for solving the SCA formulation is shorter than the time interval between two consecutive network events.

If a reconfiguration process cannot be completed before the arrival of the next network event, it is aborted and restarted after the event to follow most updated link state.

In Figure 3.1, a flow chart of the availability-aware SCR architecture for dynamic provisioning of SBPP connections is given. Firstly, an existing survivable routing algorithm, e.g., Suurballe's algorithm or successive survivable routing (SSR) [62], is chosen to provision w_c and p_c when connection request c arrives at an ingress node. The theoretical E2E availability of c which is denoted by A_c is then calculated using Equation (3.5) or (3.10) by assuming that the protection level of c is 100%. If $A_c < A_{c,SLA}$, it means that it is not possible to achieve the required availability for c with a single backup path even when its protection level is 100%, and we need to allocate another protection path or some path segments for meeting the availability constraint. In this study, the connection requests that cannot be satisfied with a single shared backup path are simply blocked.

If the connection is successfully set up, and the current network state reaches the threshold set for the next network event, the NMS will drop the current unfinished reconfiguration process and initiates another reconfiguration process by solving the SCA LP formulation. The solution tells the optimal protection level of each connection and the minimal required spare capacity along each link such that the protection level of each existing connection can be supported. In case the reconfiguration process is completed before the arrival of the next network event, the newly derived spare capacity along each link is kept in the NMS, and the source node of each existing connection is informed of the updated protection level of the connection.

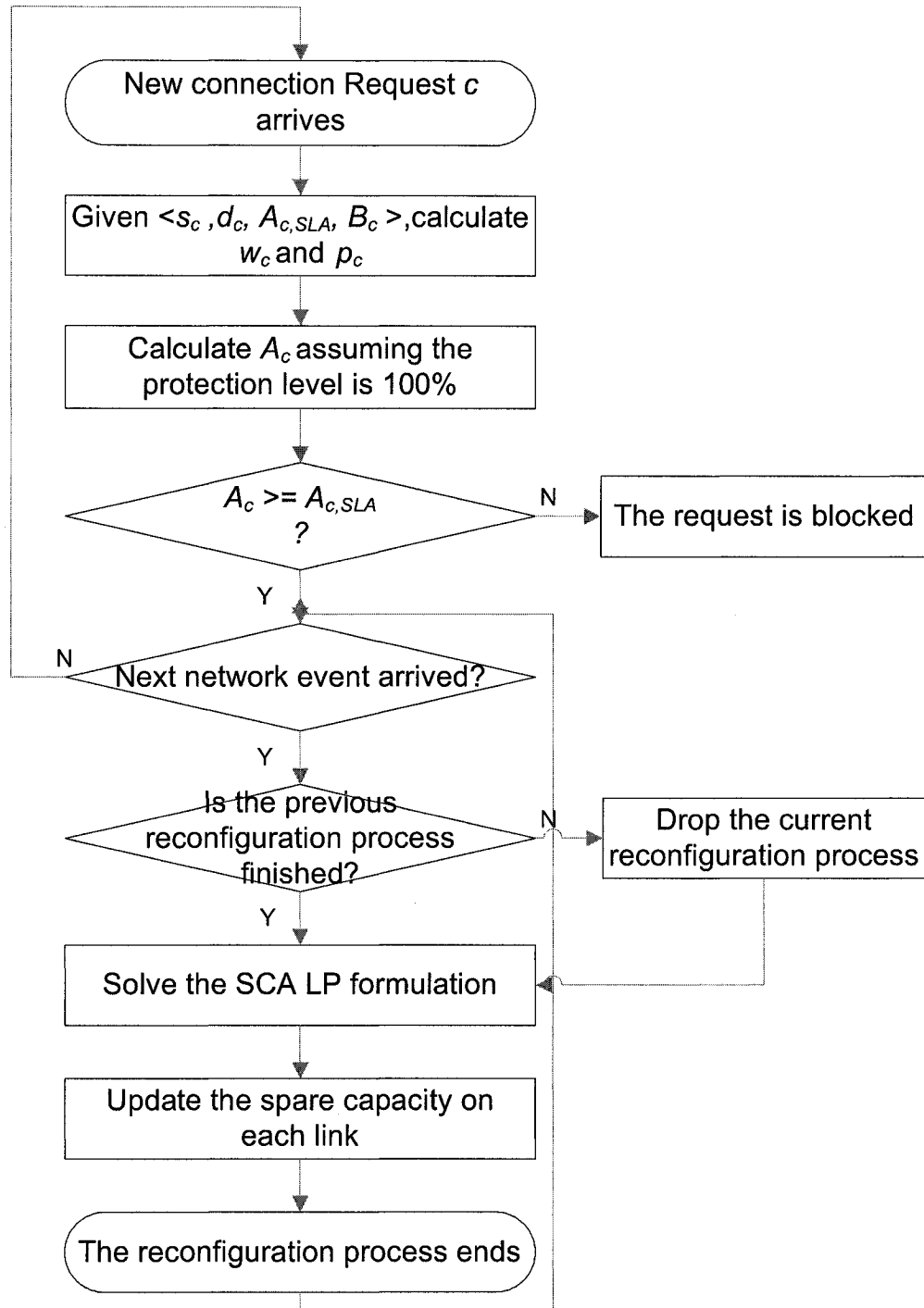


Figure 3.1: Flow chart of SCR for dynamic provisioning of SBPP connections with availability constraints

3.7 Summary

In this chapter, we have presented two policy-based mathematical models for evaluating the end-to-end availability of a shared backup path protected (SBPP) connection by assuming no more than two simultaneous failures could possibly occur in the network. With the availability-constraint of each SBPP connection, two linear programs have been formulated to optimize the spare capacity allocation (SCA) in the network. To minimize the redundancy while meeting the E2E availability requirements of SBPP connections, a new parameter is defined for each connection, called protection level, which creates a framework of partial restoration from any unexpected failure. Based on the proposed availability models and spare capacity allocation formulations, we have then presented a novel availability-aware spare capacity reconfiguration (SCR) architecture which can provision dynamic SBPP connections with differentiated protection levels according to connections' availability requirements such that the spare capacity allocated in the network is minimized.

In the next chapter, we present the numerical examples and results to verify the proposed availability models and demonstrate the efficiency of the proposed SCR architecture.

Chapter 4

Simulation Results and Analysis

4.1 Introduction

In the previous chapter, we have presented two policy-based mathematical models for evaluating the end-to-end availability of a SBPP connection. Based on the availability models, an availability-aware spare capacity allocation and reconfiguration architecture has been proposed to minimize the resource redundancy in the network. The objectives of the numerical study in this chapter are to validate our proposed availability models for SBPP connections and demonstrate the efficiency of our availability-aware SCR architecture for dynamic provisioning SBPP connections. In the following sections, we first provide detailed descriptions of our simulation models, followed by various simulation results and performance analysis.

4.2 Availability Model Validation

In this section, we verify the availability models Equation (3.5) and Equation (3.10) for SBPP connections via simulation.

4.2.1 Simulation Setup

This simulation is conducted on two sample topologies shown in Figure 4.1 and 4.2: US network and pan-European network [63] which have different network parameters. Table 4.1 shows the characteristics of each topology.

To validate the proposed availability models, a continuous time discrete event simulator is designed. Our approach is to evaluate the E2E availability of a group of connections in the network in presence of the arrival and departure of random failure events. For this purpose, firstly a set of 135 SBPP connections are randomly generated and uniformly allocated among all node pairs using successive survivable routing (SSR) algorithm [62]. All the connections are shared backup path protected.

The basic simulation assumptions are as follows:

- The link failure rate is proportional to the link length. The average failure rate $1/\text{MTTF}$ per kilometer is normalized in the unit of FIT which stands for “Failure in 10^9 hours”.
- To simulate the most realistic situation, each failure occurs independently on a link following a Poisson process, where no restriction on the number of simultaneous

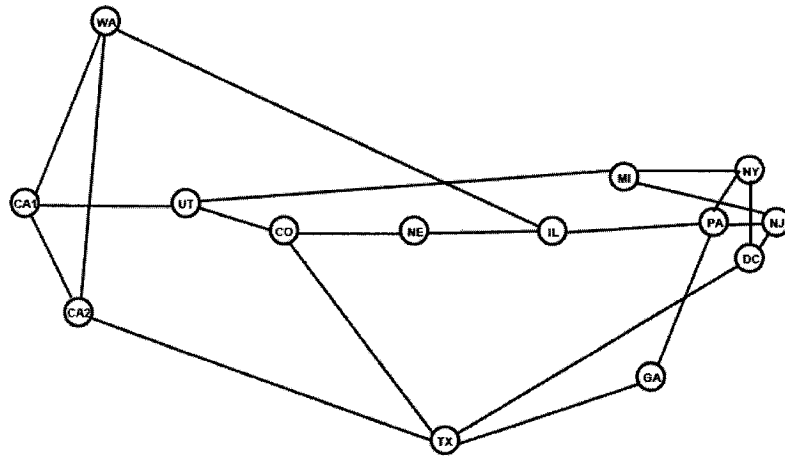


Figure 4.1: US network topology

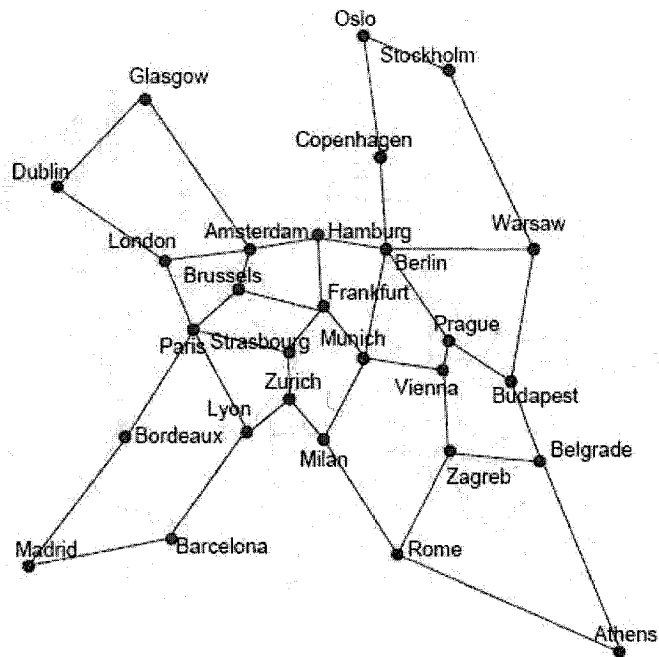


Figure 4.2: pan-European network topology

	pan-European network	US network
number of nodes n	28	14
number of links k	41	21
minimum node degree	2	2
maximum node degree	5	4
average node degree	2.92857	3
minimum link length (km)	218	312
maximum link length (km)	1500	3408
average link length (km)	625.366	1299.05
network diameter (km)	5051	5316
average distance (km)	1983.06	2722.44
network diameter (hops)	8	3
average distance (hops)	3.56085	2.14286

Table 4.1: Topological parameters of the studied networks

failures has been addressed.

- Failure holding time of each link follows a negative exponential distribution with a mean value of 6 hours, i.e., MTTR = 6 hours.
- The MTTR is assumed to be the same value on all links.
- Connections are held until the simulation is terminated.

The continuous time discrete-event simulator has been developed using C++ and executed on a LINUX server with two 2.8-GHz CPUs and 1GB memory.

4.2.2 Results from Failure-Independent Availability Model

Firstly, we compare the simulated and theoretical availabilities for each connection in the US network and pan-European network, respectively. In Figure 4.3 and Figure 4.4, we

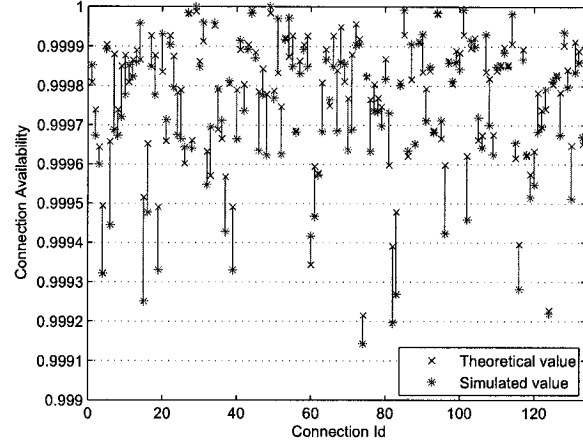


Figure 4.3: Comparison of FID-based theoretical and simulated connection availabilities for each connection when $1/\text{MTTF} = 600$ FIT and protection level is 100% in US network

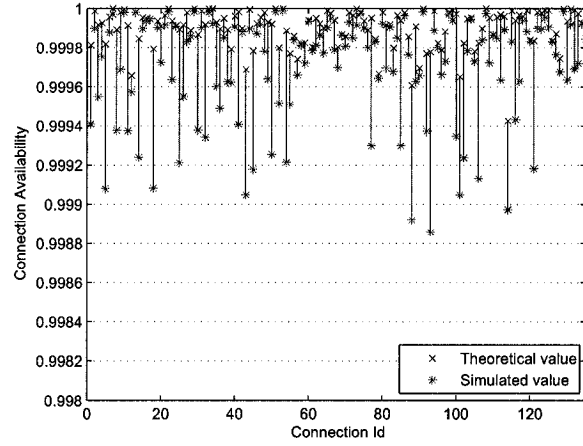


Figure 4.4: Comparison of FID-based theoretical and simulated connection availabilities for each connection when $1/\text{MTTF} = 600$ FIT and protection level is 100% in pan-European network

show the results when the protection level of all the connections is 100% and the average link failure rate $1/\text{MTTF} = 600$ FIT per kilometer. The theoretical values are computed

according to Equation (3.5). It is observed that in both networks, the theoretical E2E availability is no less than the corresponding simulated value for all the connections due to the fact that at most two links could fail simultaneously in the theoretical availability model, whereas the simulation model allowed any number of concurrent link failures to occur. Thus, a higher E2E availability could be derived than that of the realistic one. In the US network (see Figure 4.1), the difference between the simulated and theoretical availabilities (denoted as *Error*) is only about 0.00668680%, which is averaged over all connections, and the 95% confidence interval for *Error* is in the range from 0.00668600% to 0.00668760%. In the pan-European network (see Figure 4.2), *Error* is about 0.02159244%, and the 95% confidence interval for *Error* is in the range from 0.02158227% and 0.02160261%. The confidence intervals are obtained using the method described in Appendix B. These results indicate very good accuracy of the FID-based theoretical model.

Figure 4.5 and Figure 4.6 show the comparison between the simulated and theoretical availability for different protection level when $1/\text{MTTF} = 600$ FIT in the US network and the pan-European network respectively. We consider the range of protection level of each connection from 0.1 to 1 such that the lowest availability is around 99%. All the connections are assigned with the same protection level (i.e., 0.1, 0.2, 0.4, 0.6, 0.8, and 1.0) in each simulation run. Both the simulated and theoretical availability at different protection level are averaged over all the connections. In both networks, we observe that the theoretical and simulated availabilities increase along with the protection levels of the

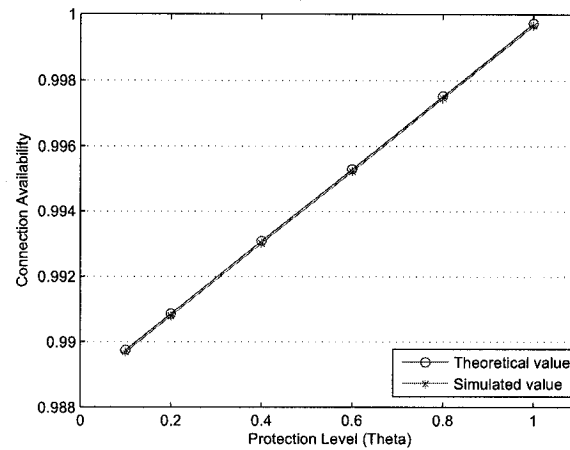


Figure 4.5: Comparison of FID-based theoretical and simulated connection availabilities at different protection levels when $1/\text{MTTF} = 600$ FIT in US network

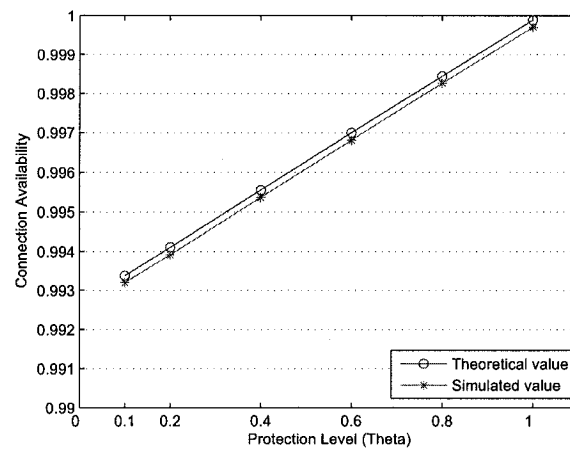


Figure 4.6: Comparison of FID-based theoretical and simulated connection availabilities at different protection levels when $1/\text{MTTF} = 600$ FIT in pan-European Network

connections. The reason is very straightforward, since the higher the protection level of a connection is, the more spare capacity it gets in case of failures. The difference between the simulated and the theoretical values changes very slightly at different protection levels

which is not large enough to be observed. In the US network, we have found that the difference between the simulated and the theoretical values is as small as 0.0066868% that can hardly be observed in the figure, while in the pan-European network, the difference is greater than that in the US network. The reason is as follows. The network diameter in terms of hops in the pan-European network is greater than that in the US network, thus, the averaged number of links that the working and protection paths of each connection traverse through is higher than that in the US network too. The probability that a connection is affected by a failure pattern with more than two link failures, which is ignored in the theoretical model, increases. Therefore, a slightly higher difference between the simulated and theoretical values is obtained in the pan-European network. Even in this case, the difference is only about 0.018492%. This validates our FID availability model for different protection levels.

We also studied the impacts of average failure rate $1/\text{MTTF}$ per kilometer on the achieved connection availabilities. Figure 4.7 and Figure 4.8 show the simulated and the theoretical connection availabilities for different $1/\text{MTTF}$ values with 100% protection level in the US network and the pan-European network respectively. Both the theoretical and the simulated values are averaged over all the connections. We observe that the curves obtained from both networks show the same trend. As we expected, the connection availability decreases when the average failure rate increases. We can also see that the difference between the theoretical and simulated availabilities is small when $1/\text{MTTF}$ is low and increases when $1/\text{MTTF}$ increases. The reason is that as the average failure

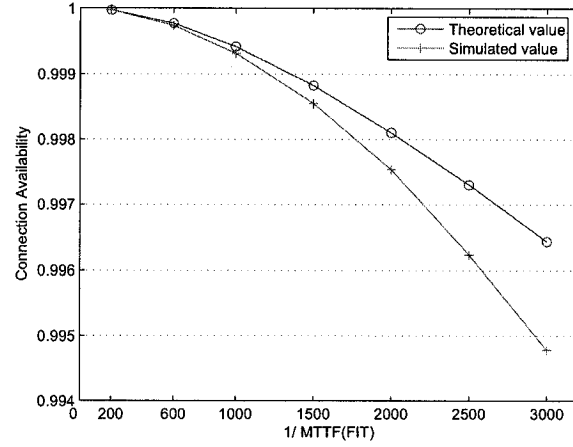


Figure 4.7: Comparison of FID-based theoretical and simulated connection availabilities at different $1/\text{MTTF}$ values when protection level is 100% in US network

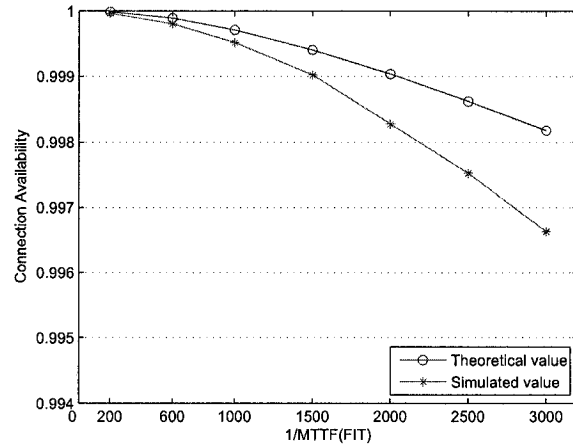


Figure 4.8: Comparison of FID-based theoretical and simulated connection availabilities at different $1/\text{MTTF}$ values when protection level is 100% in pan-European network

rate increases, there are more failure events that involves more than two links in the simulation. In the US network, the smallest difference is about 0.001597% with $1/\text{MTTF} = 200$ FIT. When the theoretical connection availability drops to 99.6437% with $1/\text{MTTF}$

= 3000 FIT, the difference is only 0.16656%. In the pan-European network, the biggest difference also occurs when $1/\text{MTTF} = 3000$ FIT and is 0.17432% which indicates a high accuracy of our FID-based availability model.

4.2.3 Results from Failure-Dependent Availability Model

In Figure 4.9 and Figure 4.10, we compare the simulated and theoretical availabilities derived from FD availability model when $1/\text{MTTF} = 600$ FIT in the US network and pan-European network respectively. For each connection, the protection levels are uniformly distributed between 0.5 and 1 for different failure patterns. The range of protection level is chosen such that the theoretical connection availability is not less than 99%, which is generally considered as the lowest availability value that the customers can accept. For failure patterns that include single or dual failures, the same protection level, which is randomly generated, is applied in the calculation of both theoretical and simulated connection availabilities. For multiple failure patterns including triple or more failures which only occur in the simulation process, the highest protection level value among that of all the dual failure patterns is adopted. It can be observed that in both networks the theoretical values are consistently higher or equal to the simulated values, which is the case similar to the FID-based availability model because only single and dual failures are considered in the theoretical model. In the US network, *Error* is about 0.01480427%, and the 95% confidence interval for *Error* is in the range from 0.01480057% to 0.01480798%. In the pan-European network, *Error* is about 0.15604744%, and the 95% confidence

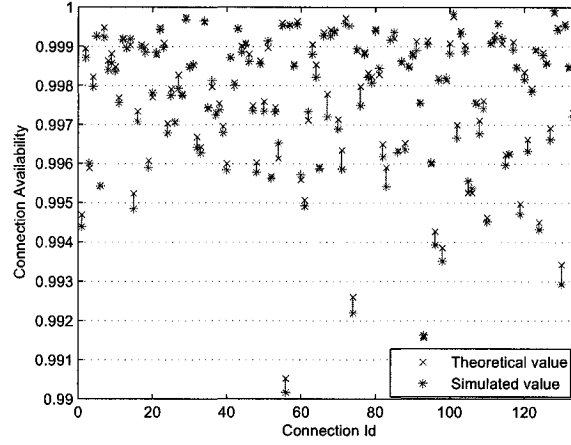


Figure 4.9: Comparison of FD-based theoretical and simulated connection availabilities when $1/\text{MTTF}=600$ FIT in US network

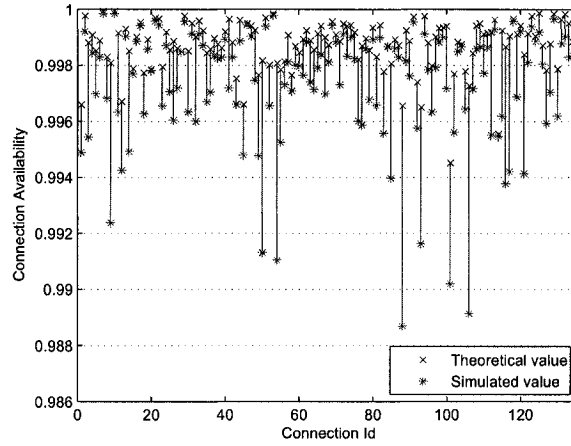


Figure 4.10: Comparison of FD-based theoretical and simulated connection availabilities when $1/\text{MTTF}=600$ FIT in pan-European network

interval for *Error* is in the range from 0.15550940% and 0.15658547%. The confidence intervals are obtained using the method described in Appendix B. These results indicate very good accuracy of the FD-based theoretical model.

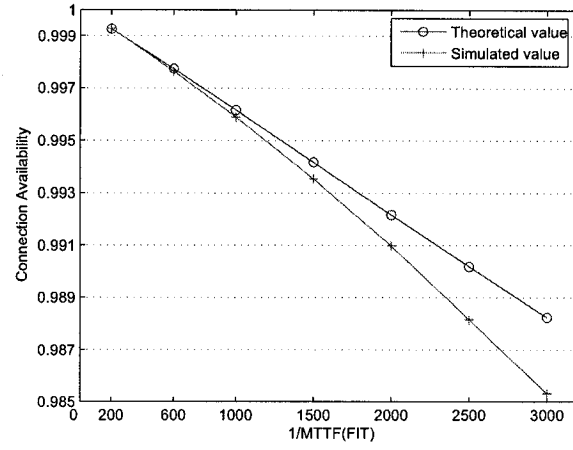


Figure 4.11: Comparison of FD-based theoretical and simulated connection availabilities at different $1/\text{MTTF}$ values in US network

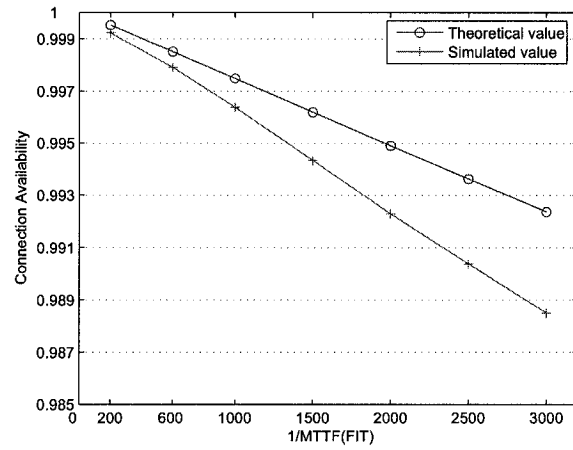


Figure 4.12: Comparison of FD-based theoretical and simulated connection availabilities at different $1/\text{MTTF}$ values in pan-European network

The comparison between the simulated and theoretical connection availabilities for different $1/\text{MTTF}$ values is shown in Figure 4.11 and Figure 4.12 for the two studied networks. Both the theoretical and the simulated values are averaged over all the con-

nections. The curves in both figures show the same trend. Similar to the case of the FID model, the difference between the theoretical and simulated values is small when $1/\text{MTTF}$ is low and increases when $1/\text{MTTF}$ increases. We can observe that the difference in the pan-European network is slightly greater than that in the US network. The reason is that the larger network diameter makes most of the connections traverse through more links in the first network, which increases the probability that a connection is affected by triple or more failures. In the US network, the smallest difference is derived when $1/\text{MTTF} = 200$ FIT with a difference of about 0.0085%. When $1/\text{MTTF} = 3000$ FIT, the theoretical connection availability is as low as 98.8227% and the difference is increased to 0.2964%. In the pan-European network, the smallest difference between the theoretical and simulated values is 0.0097% when $1/\text{MTTF} = 200$ FIT and the largest difference is 0.3943% when $1/\text{MTTF} = 3000$ FIT. These results indicate that the accuracy of the FD availability model is not as good as that of the FID availability model when the failure rate is very high. This is because the failure dependent protection levels on triple or more simultaneous link failures in the simulation are allowed and the highest protection level value among that of all the dual failure patterns is adopted.

4.3 Performance of Spare Capacity Reconfiguration Architecture

In this section, we examine the performance of our proposed SCR architecture. The two policy-based SCR schemes are compared with SSR [62], a fast and efficient SBPP connection provisioning algorithm without any availability consideration.

4.3.1 Simulation Setup

The US and pan-European network topologies shown in Figure 4.1 and Figure 4.2 are adopted in this simulation. We assume that the link failure rate depends only on the link length. The average failure rate $1/\text{MTTF}$ per kilometer is 200 FIT and the MTTR is 6 hours for all links in the network. Without loss of generality, we assume that there is no limitation on the link capacity, and each connection demands for a single unit of bandwidth and is equipped with a specific availability requirement in the range of $[0.99, 0.999999]$.

Two types of random event are defined in the network, *connection request arrival* and *connection departure*. Arrivals to k -th node-pair follow a Poisson process with an arrival rate λ_k , which is uniformly distributed in the range $[0.5\bar{\lambda}, 1.5\bar{\lambda}]$, where $\bar{\lambda}$ is the average of all the arrival rates. Each established connection has a mean holding time of $1/\mu$ with a negative-exponential distribution. We normalize time measurements using $1/\mu = 1$ so that the average traffic load between each node-pair can be considered in a

unit of Erlang as $\bar{\lambda}$.

SSR is taken to dynamically allocate each connection request. The policy-based SCA LP formulations are solved once per ten connection arrival or departure events. We are interested in the performance improvement due to the spare capacity reconfiguration. We take *spare capacity saving ratio* as the performance measure, which is defined as ratio between the amount of spare capacity saved due to the spare capacity reconfiguration and the total amount of the spare capacity consumed by SSR. The spare capacity saving ratio is measured whenever a reconfiguration process is done, and the average saving ratio is derived by averaging a certain number of saving ratios through the simulation.

The simulation program is developed using C++ and executed on a LINUX server with two 2.8-GHz CPUs and 1GB memory. The LP formulations are solved using CPLEX9.0 optimization packages [64].

4.3.2 Numerical Results and Analysis

Figure 4.13 and Figure 4.14 show the spare capacity saving ratios of FID-SCR and FD-SCR schemes in the two sample network topologies respectively. Each data point is the average saving ratio over 50 reconfiguration processes. In the pan-European network, it is clear that when the availability requirement is in a low range between 0.99 and 0.9943, the average saving ratio is very close to 100% for both schemes. This is because a protection path with a small amount of spare capacity or no spare capacity reserved at all could be enough to meet the availability constraint for most connections. With

the increase of the availability requirement on each connection, the spare capacity saving ratios of both schemes decrease since the required spare capacity on the protection path increases. When the availability requirement is above 0.9999, the saving ratio drops to less than 1%. This is because the high availability requirement on each connection can only be met with a protection level close to 100%, which is the same as that in SSR algorithm.

In the US network, the variation of the saving ratio against the availability requirement for both schemes is similar to that in the pan-European network. Some differences are made in that the saving ratios in pan-European network are always higher than those in the US network under the same availability requirements. This is because there are some very long links in the US network whose unavailabilities are relatively higher than those of links in the pan-European network; thus, when a working path traverses through these links, more spare capacity is needed on its protection path to meet the same availability requirement. We can also observe from Figure 4.13 and Figure 4.14 that the availability requirement above 0.9999 can hardly be achieved with a single protection path in the US network, while in the pan-European network, it is possible to get an availability of 0.99999.

From Figure 4.13 and Figure 4.14, it can also be observed that when the availability requirement is between 0.9943 and 0.99968 in the US network and between 0.9968 and 0.999968 in the pan-European network, the FD-SCR scheme gains an obvious advantage over the FID-SCR scheme. As discussed in Section 3.5.1, the obtained availability is

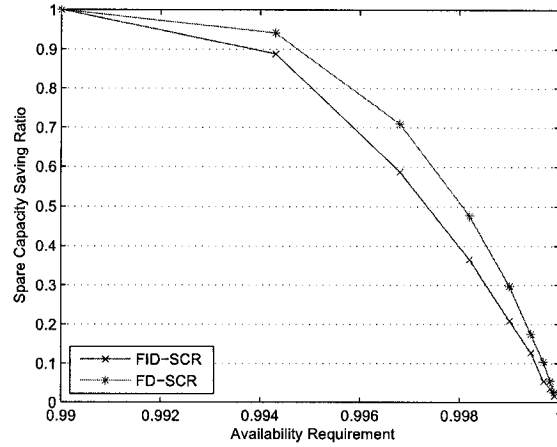


Figure 4.13: Spare capacity saving ratio vs. availability requirement in US network

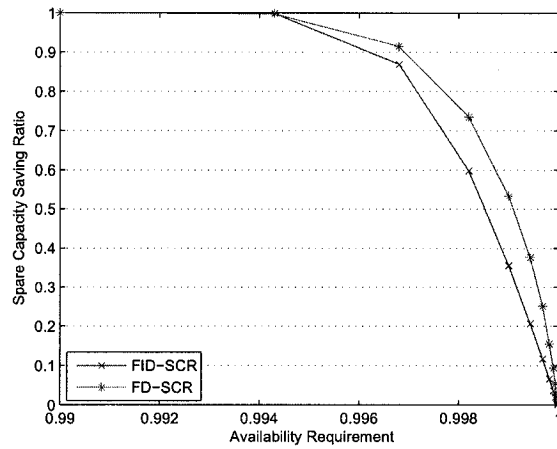


Figure 4.14: Spare capacity saving ratio vs. availability requirement in pan-European network

overestimated in the FID-SCR scheme such that less spare capacity than actually required is consumed. This leads to the fact that the derived spare capacity saving ratio with FID-SCR is generally higher than the actual saving ratio. On the other hand, even in front of the advantage taken by FID-SCR, the FD-SCR scheme still can achieve a higher

saving ratio. This result leads to a fact that failure-dependent policy can achieve better bandwidth efficiency than the failure-independent policy by making use of the location information of the failures. However, the better performance in the FD-SCR scheme is achieved at the expense of a longer computation time since the number of constraints in the FD-SCR scheme is much more than that in the FID-SCR scheme as we described in Section 3.5.2. Thus, the FID-SCR scheme can provide an actual solution although it may consume slightly more spare capacity than that of the FD-SCR scheme.

Figure 4.15 and Figure 4.16 show the protection level averaged over all the connections when different availability requirements are addressed in both sample networks. It can be observed in both figures that the protection level increases as the availability requirement increases. These results match our expectation as well as the results in Figure 4.13 and Figure 4.14 which show that the spare capacity saving ratios decrease with the increase of the availability requirements. When the availability requirement is higher than 0.9999 in the US network and 0.999968 in the pan-European network, the protection level is close to 100%. In the event that customers are interested in the availabilities in the range from 0.9968 to 0.999968 (two nine to four nine), the proposed policy-based SCR schemes for dynamic provisioning of SBPP connections will be able to provide much better resource utilization than that of the SSR algorithm.

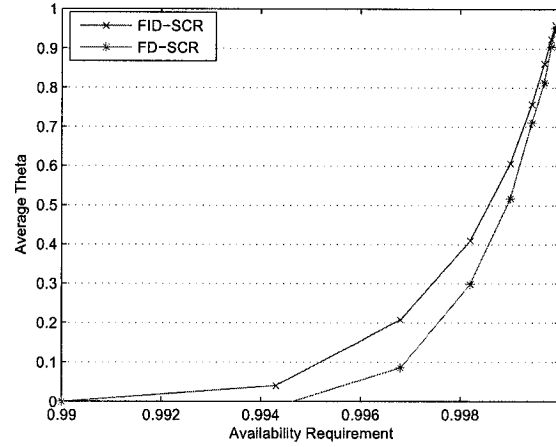


Figure 4.15: Protection level vs. availability requirement in US network

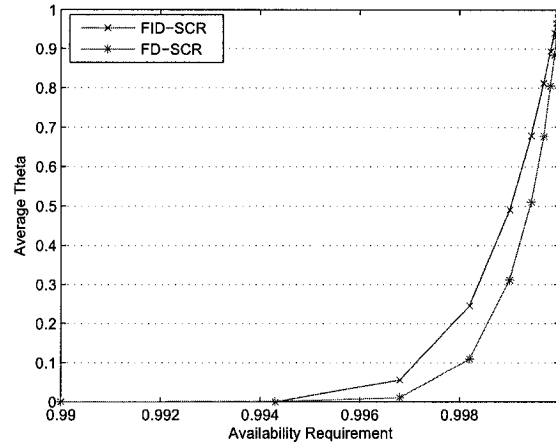


Figure 4.16: Protection level vs. availability requirement in pan-European network

4.4 Summary

In this chapter, we have provided detailed simulation results to validate our proposed availability models for SBPP connections and evaluate the performance of our availability-aware SCR architecture for dynamic provisioning SBPP connections by making perfor-

mance comparison with the SSR algorithm, which does not consider any availability constraint.

We have found that the policy-based availability models for SBPP connections have a high accuracy compared with the simulated availability values. The proposed SCR architecture can considerably reduce the spare capacity consumption. We have also found that a failure-independent policy which requires much fewer decision variables and avoids the overhead of locating the failures can provide a practical solution despite that it may require a slightly more spare capacity than a failure-dependent policy.

Chapter 5

Conclusions and Future Work

5.1 Conclusions

In this thesis, we have presented two policy-based availability models for shared backup path protected connections when no more than two simultaneous failures occur in the network. The policies could be failure-dependent or failure-independent. We have proposed to use partial protection on the shared backup paths due to the characteristics of bandwidth divisibility in the IP/MPLS layer in order to optimize the resource consumption. Based on the proposed availability models, a novel spare capacity reconfiguration architecture with availability constraints has been developed to minimize the spare capacity allocation for dynamic provisioning of shared backup path protected connections.

We have conducted a comprehensive study to verify the proposed availability models under different policies. The numerical results have shown the correctness of our availabil-

ity models. We have compared the proposed spare capacity reconfiguration architecture with SSR, a fast and efficient dynamic provisioning algorithm for shared-backup-path-protected connections without considering any availability requirement of connections. We have found that the availability-aware spare capacity reconfiguration architecture can considerably reduce the spare capacity consumption. We have also found that a failure-independent policy which requires much fewer decision variables and avoids the overhead of locating the failures can provide a practical solution despite the fact that it may require a slightly more spare capacity than a failure-dependent policy.

5.2 Future Work

As an extension to this work, the following list of open problems are suggested as potential research areas:

In the availability models for SBPP connections, we have assumed the availability impairment upon a connection due to spare capacity contention, which is caused by multiple simultaneous failures in the network, to be the worst case in order to make the mathematical formulations of spare capacity allocation linear. Thus, an upper bound on the consumed spare capacity is obtained. To further optimize the resource allocation in a network, the probability for a connection to obtain spare capacity in the presence of spare capacity contention needs to be addressed more precisely. In this work, we have also assumed no more than two simultaneous failures could possibly occur in the network. Although neglecting triple or more failure patterns only slightly overestimates

the connection availability while simplifying the availability models, to generalize the availability models, triple or more failure patterns need to be considered. Under this circumstance, accurately evaluating the availability for SBPP connections becomes extremely complicated and is an open problem to be solved.

For the spare capacity reconfiguration architecture, we have adjusted the spare capacity on each link by solving an LP formulation periodically. There are instances where the LP approaches may have difficulty due to large network size and high volume of connection requests. Developing an efficient heuristic-based connection provisioning approach, in which an appropriate level of protection is provided to each connection according to its predefined availability requirement, is a very promising topic.

Bibliography

- [1] M. Clouqueur and W. D. Grover, "Availability analysis of span-restorable mesh networks," *IEEE Journal on Selected Areas in Communications*, vol. 20, pp. 810–821, May 2002.
- [2] P.-H. Ho, J. Tapolcai, H. T. Mouftah, and C.-H. Yeh, "Linear formulation for path shared protection," in *Proc. IEEE International Conference on Communications (ICC2004)*, vol. 3, pp. 1622–1627, June 2004.
- [3] G. Li, D. Wang, C. Kalmanek, and R. Doverspike, "Efficient distributed path selection for shared restoration connections," *IEEE/ACM Trans. on Networking*, vol. 11, pp. 761–771, Oct. 2003.
- [4] S. Kini, M. Kodialam, S. Sengupta, and C. Villamizar, "Shared backup label switched path restoration." IETF draft, draft-kini-restoration-shared-backup-01.txt, May 2001. Work in progress.
- [5] D. Papadimitriou and E. Mannie, "Analysis of Generalized Multi-Protocol Label Switching(GMPLS)-based recovery mechanisms (including protection and restora-

- tion),” *IETF RFC 4428*, Mar. 2006.
- [6] P. H. Ho and H. T. Mouftah, “A framework for service-guaranteed shared protection in WDM mesh networks,” *IEEE Communications Magazine*, vol. 40, pp. 97–103, Feb. 2002.
- [7] P. Ho, J. Tapolcai, and T. Cinkler, “Segment shared protection in mesh communication networks with bandwidth guaranteed tunnels,” *IEEE/ACM Trans. on Networking*, vol. 12, pp. 1105–1118, Dec. 2004.
- [8] D. Xu, Y. Xiong, and C. Qiao, “Novel algorithms for shared segment protection,” *IEEE Journal on Selected Areas in Communications*, vol. 21, pp. 1320–1331, Oct. 2003.
- [9] D. Wang, G. Li, J. Yates, and C. Kalmanek, “Efficient segment-by-segment restoration,” in *Proc. Optical Fiber Communications (OFC 2004)*, vol. 1, pp. 23–27, Feb. 2004.
- [10] J. Zhang, K. Zhu, H. Zang, and B. Mukherjee, “A new provisioning framework to provide availability-guaranteed service in WDM mesh networks,” in *Proc. IEEE International Conference on Communications (ICC 2003), Seattle, USA*, vol. 2, pp. 1484–1488, May 2003.
- [11] L. Song, J. Zhang, and B. Mukherjee, “Dynamic provisioning with reliability guarantee and resource optimization for differentiated services in WDM mesh networks,” in

- Proc. Optical Fiber Communication Conference (OFC 2005), Anaheim, USA*, vol. 3, Mar. 2005.
- [12] Y. Huang, W. Wen, J. Heritage, and B. Mukherjee, "A generalized protection framework using a new link state availability model for reliable optical networks," *IEEE/OSA Journal of Lightwave Technology*, vol. 22, pp. 2536–2547, Nov. 2004.
- [13] M. Tornatore, G. Maier, and A. Pattavina, "Availability design of optical transport networks," *IEEE Journal on Selected Areas in Communications*, vol. 23, pp. 1520–1532, Aug. 2005.
- [14] D. A. A. Mello, D. Schupke, and H. Waldman, "A matrix-based analytical approach to connection unavailability estimation in shared backup path protection," *IEEE Communications Letters*, vol. 9, pp. 844–846, Sept. 2005.
- [15] D. Mello, J. Pelegriani, R. Ribeiro, D. Schupke, and H. Waldman, "Dynamic provisioning of shared-backup path protected connections with guaranteed availability requirements," in *IEEE/CreateNet GOSP Workshop, Boston, USA*, pp. 397–404, Oct. 2005.
- [16] D. Awduche, J. Malcolm, M. O. J. Agogbua, and J. McManus, "Requirements for traffic engineering over MPLS," *IETF RFC 2702*, Sept. 1999.
- [17] E. Rosen, A. Viswanathan, and R. Callon, "Multiprotocol label switching architecture," *IETF RFC 3031*, Jan. 2001.

- [18] D. Awduche, L. Berger, D. Gan, T. Li, V. Srinivasan, and G. Swallow, "RSVP-TE: Extensions to RSVP for LSP tunnels," *IETF RFC 3209*, Dec. 2001.
- [19] L. Andersson et al., "Constraint-based LSP setup using LDP," *IETF RFC 3212*, Jan. 2002.
- [20] D. Katz, K. Kompella, and D. Yeung, "Traffic engineering (TE) extensions to OSPF version 2," *IETF RFC 3630*, Sept. 2003.
- [21] H. Smit and T. Li, "Intermediate system to intermediate system (IS-IS) extensions for traffic engineering (TE)," *IETF RFC 3784*, June 2004.
- [22] C. Xin, Y. Ye, S. Dixit, and C. Qiao, "An agent-based traffic grooming and management mechanism for IP over optical networks," in *Proc. International Conference on Computer Communications and Networks (ICCCN 2002), Miami, USA*, pp. 425–430, Oct. 2002.
- [23] B. Rajagopalan, J. Luciani, and D. Awduche, "IP over optical networks: A framework," *IETF RFC 3717*, Mar. 2004.
- [24] E. Mannie, "Generalized multi-protocol label switching (GMPLS) architecture," *IETF RFC 3945*, Oct. 2004.
- [25] M. Tatipamula, F. L. Faucheur, T. Otani, and H. Esaki, "Implementation of IPv6 services over a GMPLS-based IP/Optical network," *IEEE Communications Magazine*, vol. 43, pp. 114–122, May 2005.

- [26] L. Berger, “Generalized multi-protocol label switching (GMPLS) signaling functional description,” *IETF RFC 3471*, Jan. 2003.
- [27] L. Berger, “Generalized multi-protocol label switching (GMPLS) signaling resource reservation protocol-traffic engineering (RSVP-TE) extensions,” *IETF RFC 3473*, Jan. 2003.
- [28] P. Ashwood-Smith and L. Berger, “Generalized multi-protocol label switching (GMPLS) signaling constraint-based routed label distribution protocol (CR-LDP) extensions,” *IETF RFC 3472*, Jan. 2003.
- [29] K. Kompella and Y. Rekhter, “Intermediate system to intermediate system (IS-IS) extensions in support of generalized multi-protocol label switching (GMPLS),” *IETF RFC 4205*, Oct. 2005.
- [30] K. Kompella and Y. Rekhter, “OSPF extensions in support of generalized multi-protocol label switching (GMPLS),” *IETF RFC 4203*, Oct. 2005.
- [31] E. J. Lang, “Link management protocol (LMP),” *IETF RFC 4204*, Oct. 2005.
- [32] S. Dong, C. Phillips, and R. Friskney, “Differentiated-resilience provisioning for the wavelength-routed optical network,” *IEEE Journal of Lightwave Technology*, vol. 24, pp. 667–673, Feb. 2006.
- [33] OIF, “User network interface (UNI) 1.0 Signaling Specification,” Oct. 2001.

- [34] ITU-T Rec. G.8080/Y.1304, "Architecture for the automatically switched optical network (ASON)," Nov. 2001.
- [35] E. Oki, K. Shiimoto, D. Shimazaki, N. Yamanaka, W. Imajuku, and Y. Takigawa, "Dynamic multilayer routing schemes in GMPLS-based IP+optical networks," *IEEE Communications Magazine*, vol. 43, pp. 108 – 114, Jan. 2005.
- [36] S. Koo, G. Sahin, and S. Subramaniam, "Dynamic LSP provisioning in overlay, augmented, and peer architectures for IP/MPLS over WDM networks," in *Proc. IEEE INFOCOM 2004, Hong Kong, China*, vol. 1, pp. 514–523, Mar. 2004.
- [37] Q. Zheng and G. Mohan, "Protection approaches for dynamic traffic in IP/MPLS-over-WDM networks," *IEEE Optical Communications*, vol. 41, pp. S24–S29, May 2003.
- [38] H. Zhang and A. Durresi, "Differentiated multilayer survivability in IP/WDM networks," in *IEEE/IFIP Network Operations and Management Symposium (NOMS)*, no. 15-19, pp. 681–694, 2002.
- [39] W. Lai and D. McDysan, "Network hierarchy and multilayer survivability," *IETF RFC 3386*, Nov. 2002.
- [40] J. Lang, B. Rajagopalan, and D. Papadimitriou, "Generalized multi-protocol label switching (GMPLS) recovery functional specification," *IETF RFC 4426*, Mar. 2006.

- [41] J. Wang, L. Sahasrabuddhe, and B. Mukherjee, "Path vs. subpath vs. link restoration for fault management in IP-over-WDM networks: performance comparisons using GMPLS control signaling," *IEEE Communications Magazine*, vol. 40, pp. 80–87, Nov. 2002.
- [42] G. Shen and W. D. Grover, "Adaptive self-protecting transport networks based on p-cycles under the working capacity envelope concept." To appear in *IEEE Journal on Selected Areas in Communications*.
- [43] W. D. Grover, "The protected working capacity envelop concept: An alternate paradigm for automated service provisioning," *IEEE Communication Magazine*, vol. 42, pp. 62–69, Jan. 2004.
- [44] D. Stamatelakis and W. Grover, "IP layer restoration and network planning based on virtual protection cycles," *IEEE Journal on Selected Areas in Communications*, vol. 18, no. 10, pp. 1938–1949, 2000.
- [45] A. Autenrieth and A. Kirstädter, "Engineering end-to-end IP resilience using resilience-differentiated QoS," *IEEE Communications Magazine*, vol. 40, pp. 50–57, Jan. 2002.
- [46] M. Pickavet, P. Demeester, D. Colle, D. Staessens, B. Puype, L. Depr, and I. Lievens, "Recovery in multilayer optical networks," *IEEE Journal of Lightwave Technology*, vol. 24, pp. 122–134, Jan. 2006.

- [47] M. Clouqueur, *Availability of Service in Mesh-Restorable Transport Networks*. PhD thesis, University of Alberta, 2004.
- [48] E. E. Lewis, *Introduction to Reliability Engineering*. New York : J. Wiley, 2nd ed., 1996.
- [49] D. Arci, D. Petecchi, G. Maier, A. Pattavina, and M. Tornatore, "Availability models for protection techniques in WDM networks," in *Proc. Design of Reliable Communication Networks (DRCN 2003)*, Banff, Canada, no. 19-22, pp. 158–166, Oct. 2003.
- [50] J. Doucette, M. Clouqueur, and W. Grover, "On the availability and capacity requirements of shared backup path-protected mesh networks," *SPIE/Kluwer Optical Networks Magazine*, vol. 4, pp. 29–44, Nov./Dec. 2003.
- [51] W. Yao and B. Ramamurthy, "Survivable traffic grooming with differentiated end-to-end availability guarantees in WDM mesh network," in *13th IEEE Workshop on Local and Metropolitan Area Networks (LANMAN 2004)*, Mill Valley, USA, vol. 25-28, pp. 87–90, April 2004.
- [52] S. Verbrugge, D. Colle, P. Demeester, R. Huelsermann, and M. Jaeger, "General availability model for multilayer transport networks," in *Proc. Design of Reliable Communication Networks (DRCN 2005)*, Ischia, Italy, no. 16-19, pp. 85–92, Oct. 2005.

- [53] M. Amin, K. Ho, G. Pavlou, and M. Howarth, "Improving survivability through traffic engineering in MPLS networks," in *Proc. 10th IEEE Symposium on Computers and Communications (ISCC 2005), La Manga del Mar Menor, Cartagena, Spain*, no. 27-30, pp. 758–763, June 2005.
- [54] P. Pongpaibool and H. S. Kim, "Novel algorithms for dynamic connection provisioning with guaranteed service level agreements in IP-over-Optical networks," in *Proc. IEEE Global Telecommunications Conference (GLOBECOM'03), San Francisco, USA*, vol. 5, pp. 2643–2648, Dec. 2003.
- [55] M. Yu and B. Xie, "An analytical availability model for MPLS networks with end-to-end IP resilience," in *Proc. IEEE Pacific Rim Conference on Computer, Communications and Signal Processing*, vol. 2, pp. 820–823, Aug. 2003.
- [56] W. D. Grover, *Mesh-based Survivable Networks: Options and Strategies for Optical, MPLS, SONET and ATM Networking*. Upper Saddle River, New Jersey: Prentice Hall PTR, 2003.
- [57] O. Gerstel and G. Sasaki, "Quality of protection (QoP) : A quantitative unifying paradigm to protection service grades," *SPIE/Kluwer Optical Networks Magazine*, vol. 3, pp. 40–50, May/June 2002.
- [58] J. Fang, M. Sivakumar, A. Somani, and K. Sivalingam, "On partial protection in groomed optical WDM networks," in *Proc. International Conference on Dependable Systems and Networks (DSN 2005), Yokohama, Japan*, pp. 228–237, June 2005.

- [59] O. Gerstel and G. Sasaki, "A new protection paradigm for digital video distribution networks," in *Proc. IEEE International Conference on Communications (ICC 2006)*, Istanbul, Turkey, pp. OS4.5.1–OS4.5.6, June 2006.
- [60] J. Zhang, K. Zhu, and B. Mukherjee, "A comprehensive study on backup repositioning to remedy the effect of multiple-link failures in WDM mesh networks," in *Proc. IEEE International Conference on Communications (ICC 2004)*, Paris, France, vol. 3, pp. 1654–1658, June 2004.
- [61] P. Ho and H. T. Mouftah, "Reconfiguration of spare capacity for MPLS-based recovery in the Internet backbone networks," *IEEE/ACM Trans. on Networking*, vol. 12, pp. 73–84, Feb. 2004.
- [62] Y. Liu, "Approximating optimal spare capacity allocation by successive survivable routing," *IEEE/ACM Trans. on Networking*, vol. 13, pp. 198–211, Feb. 2005.
- [63] A. Betker et al., "Reference transport network scenarios," in *Proc. 5th ITG Workshop on Photonic Networks*, Leipzig, Germany, 2004.
- [64] ILOG CPLEX9.0, *User's Manual*. ILOG, Mountain View, CA, USA, 2003.

Appendix A

Computation of the Stationary Probabilities of Failure Patterns

Given the values of MTTF and MTTR of each link in the network, the probabilities of all failure patterns can be derived by solving a continuous time Markov chain introduced in [14].

The basic assumptions are as follows:

- The network nodes have availability equal to one.
- A link is either available or unavailable.
- All links fail independently.
- No more than two simultaneous link failures occur in the network.

- The repair time and the time to failure of a link are memoryless, exponentially distributed random processes with constant means MTTR and MTTF.

Based on the above assumptions, the state-transition diagram for the continuous time Markov chain shown in Figure A.1 can be derived. There are three columns of states in Figure A.1. The first column contains state (0), in which none of the links have failed; the second column contains state $(1 - L)$ and exactly one link has failed in these states, where L is the number of links of the network topology; and the third column contains the states (i, j) where all links except links i and j are operating, and link i failed before link j .

The following notations will be used in the mathematical derivation:

- λ_i denotes the failure rate of link i in the network, where $\lambda_i = 1/MTTF_i$.
- μ_i denotes the repair rate for link i of the network, where $\mu_i = 1/MTTR_i$.
- π_0 denotes the probability of state (0), i.e., the proportion of time during which no links in the failed state.
- π_i denotes the probability of state (i) , i.e., the proportion of time during which link i is in the failed state.
- $\pi_{i,j}$ denotes the probability of state (i, j) , i.e., the proportion of time during which links i and j are in the failed state, where link i failed before link j .
- λ_T denotes the sum of the failure rates of all links.

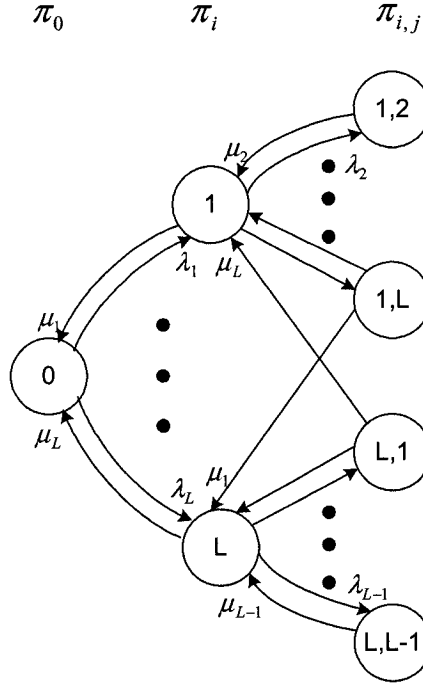


Figure A.1: The continuous time Markov chain.

The state probabilities (π) are interrelated by the balance equations:

$$(\lambda_T - \lambda_i + \mu_i) \cdot \pi_i = \lambda_i \pi_0 + \sum_{j=1, j \neq i}^L \mu_j (\pi_{i,j} + \pi_{j,i}) \quad (\text{A.1})$$

$$\pi_{i,j} = \frac{\lambda_j}{\mu_i + \mu_j} \pi_i \quad (\text{A.2})$$

$$\pi_0 + \sum_{i=1}^L \pi_i + \sum_{i=1}^L \sum_{j=1, j \neq i}^L \pi_{i,j} = 1 \quad (\text{A.3})$$

Introducing equation (A.2) into (A.1) and (A.3) yields:

$$(\lambda_T - \lambda_i + \mu_i) \cdot \pi_i = \lambda_i \pi_0 + \sum_{j=1, j \neq i}^L \left(\frac{\lambda_j \mu_j}{\mu_i + \mu_j} \pi_i + \frac{\lambda_i \mu_j}{\mu_i + \mu_j} \pi_j \right) \quad (\text{A.4})$$

$$\pi_0 + \sum_{i=1}^L \pi_i + \sum_{i=1}^L \sum_{j=1, j \neq i}^L \frac{\lambda_j}{\mu_i + \mu_j} \pi_i = 1 \quad (\text{A.5})$$

After solving (A.4) and (A.5) for π_i , (A.2) can be directly applied for calculating the stationary probability $\pi_{i,j}$ of dual-failure pattern (i, j) .

Appendix B

Confidence Interval

The accuracy of the simulation results is normally described in terms of confidence intervals. A confidence interval gives an estimated range of values which is likely to include an unknown population parameter. The estimated range can be calculated from a given set of sample data.

Let X be an unknown population parameter and $X_1, X_2, \dots, X_N$ be the simulation results of the same experiment but produced by N different runs, and assume these simulation runs are statistically independent.

The sample mean $\overline{X}$ of these results is given by

$$\overline{X} = \frac{\sum_{i=1}^N X_i}{N} \tag{B.1}$$

and the sample variance S_X^2 is defined as follows:

$$S_X^2 = \frac{\sum_{i=1}^N (X_i - \bar{X})^2}{N - 1} \quad (\text{B.2})$$

The upper and lower bounds of the confidence interval regarding these simulation results are defined as follows:

$$\text{Upper bound} = \bar{X} + \frac{S_X \times t_{\frac{\alpha}{2}, N-1}}{\sqrt{N}} \quad (\text{B.3})$$

$$\text{Lower bound} = \bar{X} - \frac{S_X \times t_{\frac{\alpha}{2}, N-1}}{\sqrt{N}} \quad (\text{B.4})$$

where $t_{\frac{\alpha}{2}, N-1}$ is the upper $100 \times \frac{\alpha}{2}$ percentage of the t -distribution with $N - 1$ degrees of freedom, and its values can be obtained from tables.

The intervals thus obtained are referred to as the intervals with $100 \times (1 - \alpha)$ percent confidence and $(N - 1)$ degrees of freedom. These confidence intervals can be made as small as desired by increasing the number of independent runs of a single experiment. In this thesis, the 95% confidence intervals for the difference between the simulated and theoretical availabilities were obtained, based on 135 independent connections.