



uOttawa

L'Université canadienne
Canada's university

**FACULTÉ DES ÉTUDES SUPÉRIEURES
ET POSTDOCTORALES**



**FACULTY OF GRADUATE AND
POSTDOCTORAL STUDIES**

Abdulbaset Taher Hassan
AUTEUR DE LA THÈSE / AUTHOR OF THESIS

M.A.Sc. (Electrical and Computer Engineering)
GRADE / DEGREE

School of Information Technology and Engineering
FACULTE, ÉCOLE, DÉPARTEMENT / FACULTY, SCHOOL, DEPARTMENT

Evaluation of Fast ReRoute Mechanisms in Broadband Networks

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

Marc St-Hilaire

Dimitrios Marrakis

Gary W. Slater

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

Evaluation of Fast ReRoute Mechanisms in Broadband Networks

Abdulbaset Taher Hassan

Thesis submitted to the
Faculty of Graduate and Postdoctoral Studies in partial fulfillment of the
requirements for the degree of

Master of Electrical and Computer Engineering
Under the auspices of Ottawa-Carleton Institute for School of
Information Technology and Engineering



University of Ottawa
Ottawa, Ontario, Canada
October 2010



Library and Archives
Canada

Published Heritage
Branch

395 Wellington Street
Ottawa ON K1A 0N4
Canada

Bibliothèque et
Archives Canada

Direction du
Patrimoine de l'édition

395, rue Wellington
Ottawa ON K1A 0N4
Canada

Your file Votre référence
ISBN: 978-0-494-74188-7
Our file Notre référence
ISBN: 978-0-494-74188-7

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 protège 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.

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

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

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


Canada

Abstract

Link failures are a common cause of service disruptions in computer networks. When a network link fails, all communications using that link are interrupted. Techniques have been deployed to improve the consequences of hardware failure networks by rerouting traffic from the failed link to other working links. The main theme of this thesis is to develop methods and tools that study and evaluate Fast ReRoute (FRR) schemes. These schemes offer fast recovery around failed components in packet networks.

First, we investigate and evaluate several network protection techniques by carrying out tests in a research testbed environment. We show the effect of link failures for several recovery scenarios, including: IGP recovery, global or end-to-end recovery, and local or segment recovery. We also study, implement and evaluate the performance of two FRR mechanisms in IP/MPLS core networks. Such mechanisms enable, in the case of a network fault, the fast switchover of protected traffic onto pre-established backup paths within 50 ms to minimize traffic loss.

In our tests, we consider single link failures protected with pre-provisioned backup paths using Traffic Engineering (TE)-FRR tunnels, or IP Fast ReRoute (IP-FRR) alternate paths. We evaluate the two techniques through experimentation. For this purpose a testbed, consisting of Juniper and Nortel equipments, is established.

Acknowledgements

All praise is due to Allah, the most merciful, the most beneficent, who has granted us knowledge and bestowed upon us good health to accomplish this work. I would also like to thank my family, especially my beloved parents, may Allah bless them and admit them to Paradise, my wife Surur and my daughter Hala. A lot of time should have been spent with you, and without your patience and support I never would have been able to complete my work.

I would like to extend my thanks and gratitude to my supervisor, Professor Hussein Mouftah of the School of Information Technology and Engineering at the University of Ottawa, for his advice, enduring patience, constant support and encouragement. Thanks are also given to Dr. Tarek Saad for his help, support and fruitful discussions. I would also like to thank my colleagues Mr. Muneer Bazama and Mr. Khaled Maamoon of the Optical Network Research Lab (ONRL) at the University of Ottawa for their productive debates and helpful ideas.

Last, but not least, I would like to express my gratitude to the staff and professors at the School of Information Technology and Engineering at the University of Ottawa for providing me with the means by which I was able to complete my degree.

Table of Contents

Abstract.....	i
Acknowledgements	ii
List of Figures.....	vi
List of Tables	viii
List of Acronyms	ix
Chapter 1	1
Introduction.....	1
1.1 Overview	1
1.2 Motivations.....	2
1.3 Thesis Objectives	3
1.4 Thesis Contributions	4
1.5 Thesis Outline	4
Related Work on Fast ReRoute Mechanisms.....	5
2.1 Synchronous Optical Network (SONET).....	5
2.2 Achieving Sub-Second IGP Convergence in Large IP Networks.....	7
2.3 Achieving Faster Failure Detection in OSPF Networks	9
2.4 Improving the Resilience in IP Networks	10
2.5 Performance Evaluation of Shortest Path Computation for IP and MPLS Multi-service Networks over Open Source Implementation	11
2.6 Performance Analysis and the Study of the behaviour of MPLS Protocols	12
2.7 MPLS Protection Switching vs. OSPF Rerouting: A Simulative Comparison.....	13
2.8 Experimental Comparison of Fault Notification and LSP Recovery Mechanisms in MPLS Operational testbeds	15
2.9 Performance Comparison of IP, MPLS and ATM Based Network Cores using OPNET ..	17
2.10 Measurement of Restoration Time from Single-Link Failure in an Optical testbed Environment.....	19
2.11 Evaluation of IP Fast ReRoute Proposals	20
Network Protection and Restoration Mechanisms	22

3.3 MPLS Two Paths Recovery	24
3.4 Local Path Protection (MPLS-TE FRR)	25
3.4.1 LDP over RSVP.....	27
3.5 IP Fast ReRoute Protection	28
3.5.1 Introduction	28
3.5.2 Recovery Mechanisms in IP-FRR	28
3.5.2.1 IP FRR using Loop-Free Alternate	28
3.5.2.1.1 Recovery Process	29
3.5.2.1.2 Pros and Cons.....	31
3.5.2.2 IP-FRR using Not-via-addresses.....	32
3.5.2.2.1 Recovery Process	33
3.5.2.2.2 Pros and Cons.....	33
3.5.2.3 IP-FRR using U-turn Alternate	34
3.5.2.3.1 Recovery Process	35
3.5.2.3.2 Pros and Cons.....	36
3.5.2.4 IP-FRR using Tunnels.....	36
3.5.2.4.1 Recovery Process	36
3.5.2.4.2 Pros and Cons.....	37
3.6 Summary	38
Chapter 4. Experiments and Analysis.....	39
4.1 Hardware used for the Experiments	39
4.2 Network Recovery Scenarios	42
4.2.1 Description of Testbed Hardware.....	42
4.2.2 Experiment Methodology	43
4.2.3 Three Network Recovery Scenarios	45
4.2.3.1 Scenario#1 (Layer-3 OSPF routing recovery)	45
4.2.3.2 Scenario#2 (MPLS two paths recovery)	46
4.2.3.3 Scenario#3 (MPLS local recovery).....	48
4.2.4 Network Recovery Result Comparison	50
4.3 Performance of Fast ReRoute Mechanisms	52
4.3.1 Scenario#4 (Number of LSPs versus Convergence time)	53

4.3.2 Scenario#5 (Number of Hops versus Convergence time)	61
4.3.3 Scenario#6 (Number of IGP Prefixes versus Convergence time)	64
4.4 Summary	67
Chapter 5. Concluding Remarks and Future Work.....	68
5.1 Concluding Remarks	68
5.2 Future Work	69
References	71
Appendix A: Router Configuration	77
Appendix B: MPLS-TE Work Verification.....	83
Appendix C: IP-FRR Work Verification.....	86
Appendix D: Confidence Interval Calculations	91

List of Figures

Figure 2.1: UPSR Self healing Ring	6
Figure 2.2: BLSR Self healing Ring	6
Figure 2.3: Convergence time for 23 Router Failures of Tier-1 ISP	8
Figure 2.4: Basic Last-hop Cell of an O2 Network	10
Figure 2.5: Comparisons of Mean Losses during a Link Failure	12
Figure 2.6a: Traditional IP Network and Figure 2.6b: MPLS Network	12
Figure 2.7: Cost239 Network.....	13
Figure 2.8: Restoration Time	14
Figure 2.9: Operational MPLS Networks	15
Figure 2.10: Network Topology for Simulation	17
Figure 2.11: End-to-end Delay and Dropped Packets	18
Figure 2.12: Optical testbed Environment.....	19
Figure 2.13: Packet Loss and Restoration Time	20
Figure 2.14: Mechanisms Comparison Node/Link Failures	21
Figure 3.1: MPLS two Paths Recovery.....	24
Figure 3.2: MPLS-TE FRR Protection	26
Figure 3.3: LDP over RSVP	27
Figure 3.4: Loop Free Alternate Scheme.....	29
Figure 3.5: Loop Free Alternate Link Failure.....	30
Figure 3.6: Link Protection causing Micro Loops	31
Figure 3.7: Not-via-addresses Scheme	32
Figure 3.8: Not-via-addresses Link Failure	33
Figure 3.9: U-turn Alternate Scheme.....	34
Figure 3.10: U-turn alternate Link Failure.....	35
Figure 3.11: Using Tunnels Link Failure.....	37
Figure 4.1: Testbed used for our Experiments.....	40
Figure 4.2: Network Topology used for Network Recovery Scenarios.....	43
Figure 4.3: Primary Link Failure in Network Recovery Scenarios	44
Figure 4.4: Traffic dropped in Layer-3 OSPF routing Recovery Scenario.....	45
Figure 4.5: Jitter in Layer-3 OSPF routing Recovery Scenario.....	46

Figure 4.6: Traffic dropped in MPLS two Paths Recovery Scenario	47
Figure 4.7: Jitter in MPLS two Paths Recovery Scenario	48
Figure 4.8: Traffic Dropped in MPLS Local Recovery Scenario	49
Figure 4.9: Jitter in MPLS Local Recovery Scenario	49
Figure 4.10: Traffic Dropped Comparison	51
Figure 4.11: Jitter Comparison	52
Figure 4.12: Link Failure in Multiple LSPs Scenario.....	53
Figure 4.13: Link Failure in one Protected LSP Scenario	54
Figure 4.14: Distribution of 10 LSPs	55
Figure 4.15: Distribution of 20 LSPs.....	56
Figure 4.16: Distribution of 40 LSPs.....	56
Figure 4.17: Distribution of 80 LSPs.....	57
Figure 4.18: Distribution of 150 LSPs.....	57
Figure 4.19: Distribution of 300 LSPs.....	58
Figure 4.20: Traffic Disruption in Multiple LSPs Scenario	59
Figure 4.21: Number of LSPs vs. Convergence time with Confidence Interval	60
Figure 4.22: Link Failure in Multiple Hops Scenario.....	61
Figure 4.23: Traffic Disruption in Multiple Hops Scenario	62
Figure 4.24: Number of Hops vs. Convergence time with Confidence Interval	63
Figure 4.25: Link Failure in IP-FRR Scenario.....	64
Figure 4.26: No. of IGP Prefixes vs. Convergence Time with Confidence Interval	65
Figure 4.27: Traffic Disruption in IP-FRR Scenario	66

List of Tables

Table 2.1: Network Topologies used in Simulations.....	9
Table 2.2: Failure Detection and Recovery Times (FDT, FRT).....	9
Table 2.3: Measurements Results	16
Table 4.1: Comparison of Different Network Recovery Schemes	50
Table D.1: Confidence Interval Calculations in Multiple LSPs Scenario	93
Table D.2: Confidence Interval Calculations in Multiple Hops Scenario	93
Table D.3: Confidence Interval Calculations in IP-FRR Scenario	94

List of Acronyms

APS	Automatic Protection Switching
ATM	Asynchronous Transfer Mode
BFD	Bidirectional Forwarding Detection
BGP	Border Gateway Protocol
BLSR	Bidirectional Line Switched Ring
CBR	Constraint Bit Rate
CE	Costumer Edge
CR- LDP	Constraint-based Routing LDP
CSPF	Constrained Shortest Path First
DWDM	Dense Wavelength Division Multiplexing
ECMP	Equal Cost Multipath Protocol
FDT	Failure Detection Time
FIB	Forward Information Base
FIS	Fault Indication Signal
FRR	Fast ReRoute
FRT	Failure Recovery Time
FTP	File Transfer Protocol
GE	Gigabit Ethernet
GPS	Global Positioning System
IETF	Internet Engineering Task Force
IGP	Interior Gateway Protocol
IP	Internet Protocol
IP-FRR	IP Fast ReRoute
IPTV	Internet Protocol Television
IS-IS	Intermediate System to Intermediate System
ISP	Internet Service Provider
LAN	Local Area Network
LDP	Label Distribution Protocol
LER	Label Edge Router
LFA	Loop Free Alternate

LSP	Label Switched Path
LSR	Label Switched Router
MP	Merge Point
MPLS	Multiprotocol Label Switching
NHOP	Next Hop
NNHOP	Next Next-Hop
NTP	Network Time Protocol
OSPF	Open Shortest Path First
PATH ERR	Path Error
PE	Provider Edge
PLR	Point of Local Repair
POPs	Point Of Presence
PSTN	Public Switched Telephone Network
QoS	Quality of Service
RSVP	Resource Reservation Protocol
SLRG	Shared Link Risk Group
SONET	Synchronous Optical Network
SP	Service Provider
SPF	Shortest Path First
TCP	Transmission Control Protocol
TE	Traffic Engineering
TE-FRR	Traffic Engineering Fast ReRoute
UDP	User Datagram Protocol
UPSR	Unidirectional Path Switched Ring
WDM	Wavelength Division Multiplexing

Chapter 1

Introduction

1.1 Overview

For most people, the Internet is simply a magical technology delivering a set of services that they rely on in their daily life, whether at home or at work. However, these applications require that all data arrives at the destination on time. Unfortunately, delay has received little attention, and the TCP-IP protocol suite has been developed in the early years of the Internet to provide communication properties required for the common applications. IP protocols have been developed to ensure packet formation, addressing and routing, while TCP has been designed to ensure successful transmission and in-order delivery of packets from end-to-end. Robust routing has been one of the design goals from the beginning. Re-convergence of the network state in case of failure has been the approach specified to ensure robust routing. This re-convergence has been perfectly compatible with the size of networks and the delay tolerant applications foreseen in the first phase [CLA88].

During the early and mid nineties, the Internet became publicly available with browsing and email being the key services. More recently, services like Voice over IP, IPTV, and online gaming have been deployed on a larger scale. These applications have considerable higher requirements to timeliness and fast failure handling during failures. The importance of these requirements is due to the real-time characteristics. The re-convergence of IP does not offer a time-scale applicable to these services.

In addition, traditional IP networks do not offer any guarantees of providing the sensitive application needs of faster recovery times in the case of failures. Therefore, many research studies have been conducted in an attempt to tackle this problem by providing some QoS guarantees for specific kinds of traffic post-failure events [MET02].

In this thesis, we provide methods and tools that have been used to evaluate the performance of two FRR network protection schemes by carrying out several experiments in a research testbed lab environment. Our testbed consists of real routers and switches from different vendors.

Moreover, we analyze results of the tests for the two schemes, as well as show improvements to IP routing in terms of providing faster recovery of traffic in failure scenarios.

The next sections specify the research motivations and objectives that are the key themes of this thesis. More detailed and specific solutions along with related works are provided in Chapters 2 and 3.

1.2 Motivations

Most transport networks are designed to recover from failures. End-to-end automatic protection switching and restoration are common approaches in circuit switched networks such as optical WDM and SONET networks that operate at the physical or Data link layer. Although these approaches can provide up to 50 ms recovery time [GER99], they require dedicated resources and hardware, which can be very expensive for service providers. For this reason, many researches have been conducted where the network layer plays the main role in the recovery process.

Until recently, recovery at the IP layer has been handled by the slow process of global re-convergence, which is described in Chapter 3. The slow IP recovery has been one of many motivations for introducing Multi-Protocol Label Switching (MPLS). The MPLS framework [ROS01] specifies mechanisms for global and local protections, including MPLS-TE FRR for local span recovery [SHA03], which is also described in more detail in Chapter 3.

Since MPLS appears to solve the problem of fast recovery in IP networks, should we then invest a lot of effort to provide the same solutions for pure connectionless IP operation?

We argue that IP-FRR is necessary for several reasons. First, not all IP networks support MPLS. This is particularly true when looking at wireless ad hoc networks, where MPLS is currently not

deployed. Furthermore, the scalability of establishing and managing the amount of MPLS backup paths required to protect all components has been questioned. Moreover, MPLS-TE using Resource Reservation Protocol (RSVP) for signalling means more control plane overhead traffic. On the other hand, IP Fast ReRoute (IP-FRR) uses little or no signalling. In terms of cost, IP-FRR requires minimal management by the operator and little or no operator training is required.

1.3 Thesis Objectives

The main objective of this thesis is to develop and implement methods and tools that investigate and evaluate the performance and scalability of FRR mechanisms in a packet core network experimentally in the lab. Moreover, these methods and tools allow us to study the different technologies that have been proposed so far, in order to provide QoS guarantees to delay sensitive applications, such as voice and video in broadband networks. In addition, such a tools and methods show the importance needs of FRR in IP networks. We then present the pros and cons of each technique in order to answer the question that many Internet service providers and companies have asked: In terms of performance, scalability, and cost, why do we need a separate IP-FRR solution, why is it not enough to use an MPLS detour to next-hop to protect links?

Using IP and MPLS routers from different vendors, we have implemented the two solutions based on IP-FRR and MPLS-TE FRR on two testbed setups at the Optical Network Research Lab (ONRL) at the University of Ottawa.

Link failure tests in the test network are simulated by pulling out fibers or shutting down one of the link between the two routers. For traffic flow, we have used the Iperf [UOF03] tool to send and receive the test packets. We have repeated our tests a number of times in order to compute the confidence intervals and ensure accuracy of our results.

1.4 Thesis Contributions

The contributions provided in this thesis can be summarized as follows:

1. The implementation of methods and tools that inspect and evaluate different Fast ReRoute protection schemes, by setting up a network using real routers and switches from different vendors.
2. The provision of interoperability solutions to insure end-to-end functionality.
3. The provision of the performance evaluation of all IP/MPLS recovery mechanisms such as, Layer-3 OSPF routing, MPLS two paths, and MPLS local recovery based on actual experimental lab measurements. As well as, a detailed comparison is given for the three different techniques.
4. The evaluation of two Fast ReRoute recovery schemes in case of failure for two different networks, i.e. pure IP networks and MPLS networks. The convergence time and packet loss have been measured in each case, and the confidence intervals have been computed. An analysis of the collected results for both cases is then presented to draw some conclusions as to the effectiveness and applicability of each scheme.
5. The provision of some important ideas and configurations tips for network operators to architect a fast recovery network environment.

1.5 Thesis Outline

The remainder of this thesis is organized as follows. Chapter 2 presents a review of the relevant background related work on FRR mechanisms. Chapter 3 discusses the protection and restoration techniques at the network layer and show the differences and the application for each technique in a scientific perspective. It presents and evaluates MPLS-TE FRR as a scheme that offers faster recovery times in MPLS core networks. It also describes the IP-FRR as another potential solution as a fast re-convergence protocol. Chapter 4 presents a description of the several experiments that are conducted over the test networks as well as analysis of the obtained results. Finally, Chapter 5 concludes this thesis and provides some ideas for future research.

Appendices A, B, and C show router configuration samples as well as some verified working command that are used in our MPLS-TE FRR and IP-FRR tests. Appendix D shows an example of the computation of the Confidence Intervals for the results that are obtained from our tests.

Chapter 2

Related Work on Fast ReRoute Mechanisms

The convergence of any network after a link or node failure has been studied by various authors and companies for many years. As a result, many research papers have been published and many tests have been conducted on either simulated environments or, as in our case, testbed environments. Moreover, there have been recent proposals regarding changes of existing protocol specification in order to improve the routers and increase the speed of the convergence process.

In this chapter, we investigate such publications as well as discuss tests that are carried out in different environments, starting with the physical layer solutions, such as SONET, then with network layer solutions, such as Layer-3 OSPF routing recovery, MPLS path recovery, and concluding with Fast ReRoute solutions that provide competitive convergence times to SONET's recovery time.

2.1 Synchronous Optical Network (SONET)

SONET technology has been universally accepted as a network technology that has provided very fast protection switching times. SONET adopts the protection schemes and is usually configured as two counter rotating rings [ANE99]. The SONET protection mechanism is called Automatic Protection Switching (APS) and its restoration time can reach up to sub-50ms, which required typically about 20 ms for fault detection, 10 ms for signalling, and 10 ms for the operation of the tail-end transfer relay; consequently, the specification for APS switching times has been reasonably set at 50 ms, allowing a 10 ms margin. APS includes the Unidirectional Path Switched Ring (UPSR) scheme and the Bidirectional Line Switched Ring (BLSR) scheme.

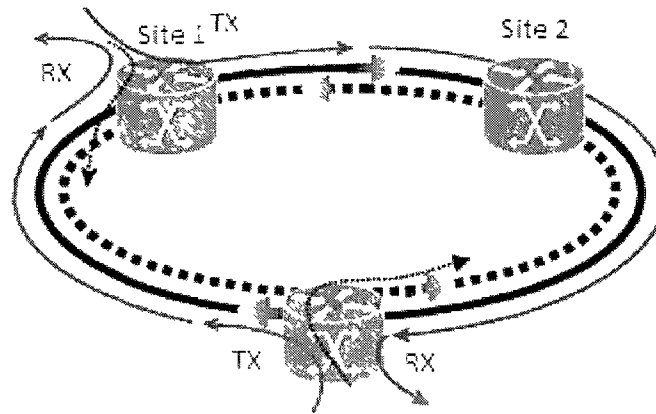


Figure 2.1: UPSR Self healing Ring [GER99]

With UPSR [GER99], the traffic is sent to both rings simultaneously, and the destination node can decide to receive the traffic on either ring as shown in Figure 2.1. In case of failure, UPSR is based on sending two redundant copies of protected traffic around a ring in two different fibres. A selector at the egress node determines the higher-quality copy in order to use the best copy.

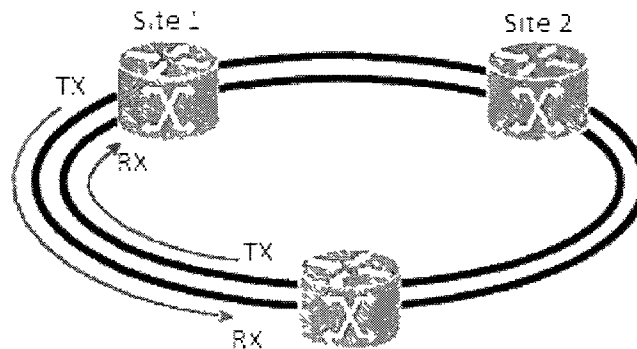


Figure 2.2: BLSR Self healing Ring [GER99]

With BLSR [GER99], the traffic is spread on both rings with no more than half of the link capacity in normal condition. The second half of the bandwidth of each ring is reserved for the protection of the other ring as shown in Figure 2.2. In case of failure, BLSR does not send redundant copies from ingress to egress. It is the ring nodes adjacent to the failure that reroute the traffic around the ring. BLSR is slightly more complex than UPSR but the bandwidth efficiency is greater.

In terms of protection, Synchronous Optical Network (SONET) is designed to switch traffic from a primary path to a secondary path in less than 50 ms. However, SONET requires dedicated hardware and resources which are very expensive. Consequently, many researchers study recovery on the network level to give the competitive convergence time that has been provided by SONET.

2.2 Achieving Sub-Second IGP Convergence in Large IP Networks

In [FRA05a] the authors describe and analyze in detail the various factors that influence the convergence time of intra domain link state routing protocols. To characterize the convergence process, the researchers first use detailed measurements in order to determine the time required to perform the various operations of a link state protocol on currently deployed routers. They then build a simulation model based on those measurements and use it to study the convergence time in large networks.

The authors in this work use two representatives but very different Service Provider (SP) topologies. The first one is GEANT, which is the pan-European Research Network. GEANT connects all the National Research networks in Europe and has interconnections with research networks in other continents. The second studied network contains the backbone nodes of a worldwide Tier-1 ISP. It is representative of a large commercial SP network.

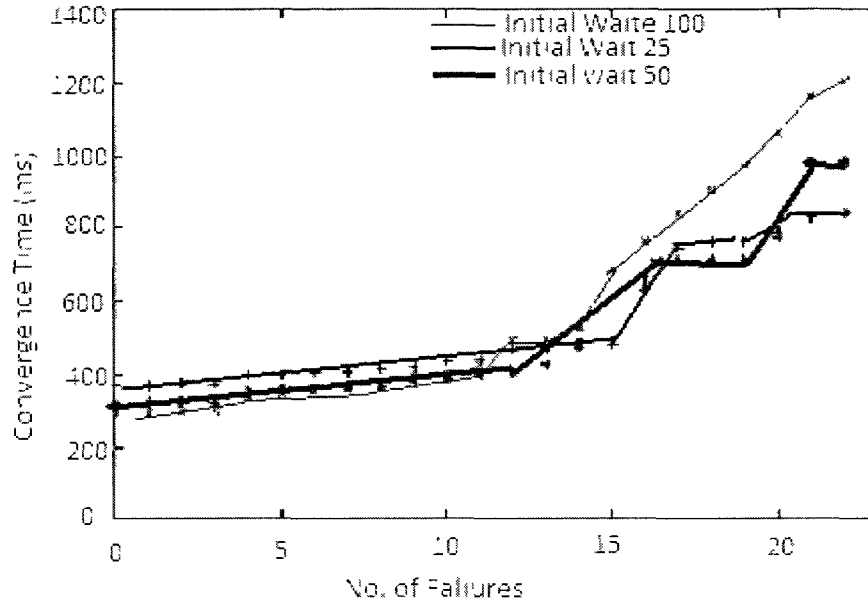


Figure 2.3: Convergence Time for 23 Router Failures of Tier-1 ISP [FRA05a]

Figure 2.3 shows the convergence time after a link failure of Tier-1 ISP network with full Forward Information Base (FIB) update, fast flooding off, and pacing 33 ms. From the simulation results, we see that in case of link failures, a convergence time of a few hundred of milliseconds can be achieved by using a low initial wait timer for the SPF computation and incremental FIB updates. Moreover, the results show that advertising fewer prefixes in the IGP significantly reduces the convergence time. When considering router or Shared Link Risk Group (SRLG) failures, convergence time is only slightly larger, provided that the pacing timer is disabled for urgent LSPs and that the initial wait timer is not too low.

Overall, the analysis shows that with current router technology sub-second IGP, convergence can be provided without any compromise on stability. However, the authors achieved the sub-sec in IGP by reducing the wait timer, which usually leads to entering unwanted recovery processes in congestion network.

2.3 Achieving Faster Failure Detection in OSPF Networks

In [GOY03] the authors attempt to find the optimal value for the Hello Interval that would lead to faster failure detection in the network, while keeping the occurrence of false alarms within acceptable limits by reducing the value of Hello Interval. Moreover, they inspect the effectiveness of this fast failure detection in achieving fast failure recovery in OSPF network.

Topology	Nodes	Links	Topology	Nodes	Links
A	9	72	D	37	88
B	27	58	E	51	176
C	27	116	F	116	476

Table 2.1: Network Topologies Used in Simulations [GOY03]

This assessment is carried out through simulations on network topologies of commercial ISPs using a detailed implementation of the OSPF protocol in NS2 simulator, which models all the protocol features as well as various standard and vendor introduced delays in the functioning of the OSPF protocol. Table 2.1 lists some of the characteristics of these topologies. Most of the topologies are irregular. However, topology A is a pure mesh while topology B has a star-like structure.

Hello Interval	Seed 1		Seed 2		Seed 3	
	FDT	FRT	FDT	FRT	FDT	FRT
10s	32.08s	36.60s	39.84s	46.37s	33.02s	38.07s
2s	7.82s	11.68s	7.63s	12.18s	7.79s	12.02s
1s	3.81s	9.02s	3.80s	8.31s	3.84s	10.11s
0.75s	2.63s	7.84s	2.97s	5.08s	2.81s	7.82s
0.5s	1.88s	6.98s	1.82s	6.89s	1.79s	6.85s
0.25s	0.95s	10.24s	0.84s	6.08s	0.99s	13.41s

Table 2.2: Failure Detection and Recovery Times (FDT, FRT) [GOY03]

From the simulation results, it is indicated that the optimal value for the Hello Interval in a network is strongly influenced by the expected congestion levels and the number of links in the

topology. While the Hello Interval can be much lower than current default value of 10s, it may not be advisable to reduce it to the millisecond range. Further, it is difficult to prescribe a single Hello Interval value that performs optimally in all cases. The network operator should set the Hello Interval conservatively, taking in account both the expected congestion levels as well as the number of links in the network topology.

2.4 Improving the Resilience in IP Networks

The authors [SCH03] propose an algorithm that uses multiple active paths at any node and towards any destination in order to protect networks against link failures. When a node locally detects a failed link or node, it immediately removes the defective element from the forwarding table and continues using the remaining next hops for forwarding packets. This technique is similar to IP Fast ReRoute mechanism but it is referred in this paper as fast local reaction, in contrast to a reaction involving other components of the network. The proposed algorithm is called an “O²” algorithm (for “out degree 2”, using graph theory terminology) as shown in Figure 2.4.

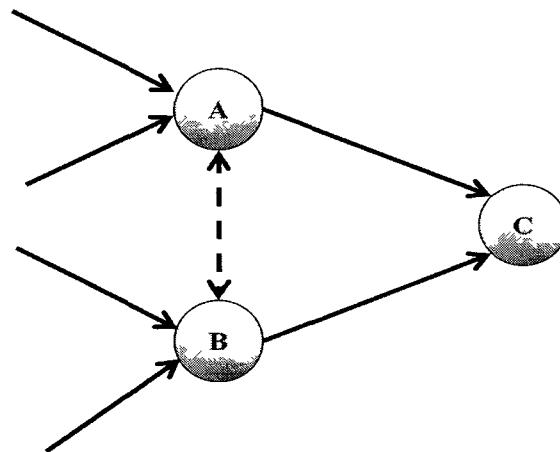


Figure 2.4: Basic last-hop Cell of an O2 Network [SCH03]

Figure 2.4 shows a basic routing in an O2 network at the last hop towards a destination C. Nodes A and B are both neighbours of C and are linked to each other. To make A and B O2 nodes, the latter link, shown as a broken line, must be used in either direction for packets towards node C if one of the direct links towards C fails. In order to prevent routing loops, the link A-B is not used for traffic towards node C unless one of the links A-C or B-C fails. Such a link is called a “joker

link” as it can be locally used when needed by any of the nodes A or B without first informing the node at the other end.

As a result, the proposed algorithm shows that it provides routes for loop free destination based multi-path routing. Although multi-path routing slightly increases the link loads in normal operation by sometimes using longer paths than necessary, overall network capacity can be saved because multi-path routing distributes the load change after a link failure in the network.

2.5 Performance Evaluation of Shortest Path Computation for IP and MPLS Multi-service Networks over Open Source Implementation

The authors [ABD04] study the performance evaluation of IP and MPLS networks in case of link failure events. Both IP and MPLS are based on the Shortest Path Computation routing paradigm. In the experiment, five different traffic flows are used in a multi-service network that consist of a Public Switched Telephone Network (PSTN), an Asymmetric Digital Subscriber Line (ADSL) access network , two local area networks (LANs), a wireless LAN and a MPLS/Diffserv core. The duration of the traffic flow is 60 seconds. The link failure occurs within this interval specifically between the seconds 10 to 20. The measures consider recovery time, latency, and packet losses.

As a result, in the IP network scenario, the failure occurs past the 15 second mark and takes almost 10 seconds to recover the routes and retransmit the data. In the MPLS network scenario, the link failure also occurs during the same time interval, but the recovery time of the MPLS network is shorter than that of the IP, and there is less packet loss.

In terms of losses, with respects to IP networks, the period of loss is greater, and the traffic flows also suffer losses. It can also be observed that after the link failure, the behaviour of losses remains the same. With respects to the MPLS platform, 100% losses occur only for CBR1 and CBR2 flows in the failure interval. Also, the losses appear to increase after the failure interval.

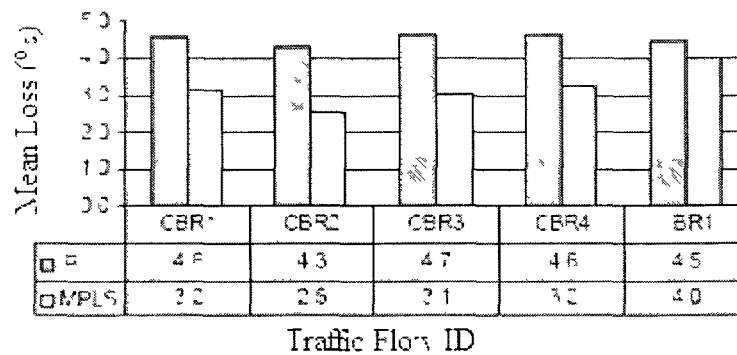


Figure 2.5: Comparisons of Mean Losses during a Link Failure [ABD04]

To conclude, Figure 2.5 shows a comparison of mean losses of both platforms. For every flow, the MPLS platform shows a better performance in terms of loss percentage. In this work, we support this finding by conducting tests on the two platforms (IP and MPLS).

2.6 Performance Analysis and the Study of the behaviour of MPLS Protocols

In [RAH08], a full comprehensive simulation environment is developed to underline the impact of Traffic Engineering over the traditional network and to evaluate the comparative performance of network traffic behaviour.

Based on the comparison of signalling protocols, it can be noted that RSVP has a drawback in its scalability, for this reason choosing CR-LDP for MPLS protocol is favourable.

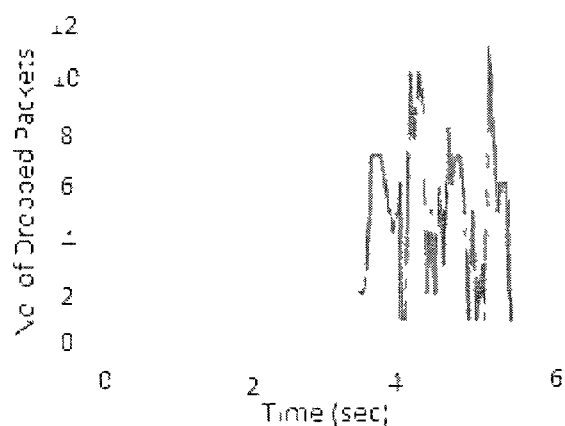


Figure 2.6a: Traditional IP Network [RAH08]

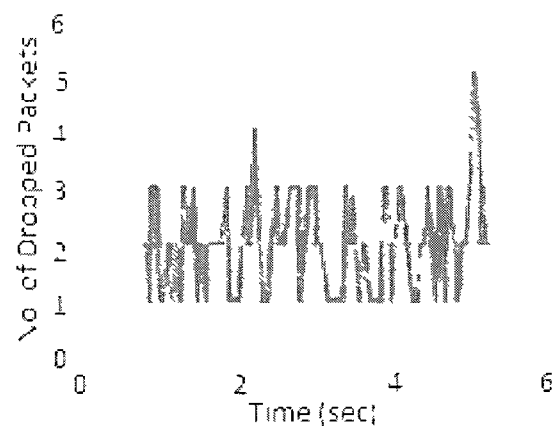


Figure 2.6b: MPLS Network [RAH08]

As shown in Figures 2.6a and 2.6b above, the packet loss is less in MPLS networks than in that of traditional networks. Furthermore, packet transmissions (in terms of both delay and loss) are improved in MPLS networks. Although the chosen parameters can be disputed to be artificially extreme, the traffic engineering mechanism improves the performance of general delay and loss. However, this work is carried out in a simulated environment, which is not accurate enough to evaluate the performance of both traditional and MPLS networks. In our study, we support these results by using a real network including routers and switches from various vendors.

2.7 MPLS Protection Switching vs. OSPF Rerouting: A Simulative Comparison

Pasqualini et al. [PAS08] investigate and compare, by implementing a simulation tool, the recovery performances of OSPF and MPLS. In addition, their focus has been on the speed of the traffic restoration after failure.

The Pan-European optical network from the COST 239 project is chosen as the main sample network because of its widespread use for network investigations (Figure 2.7).

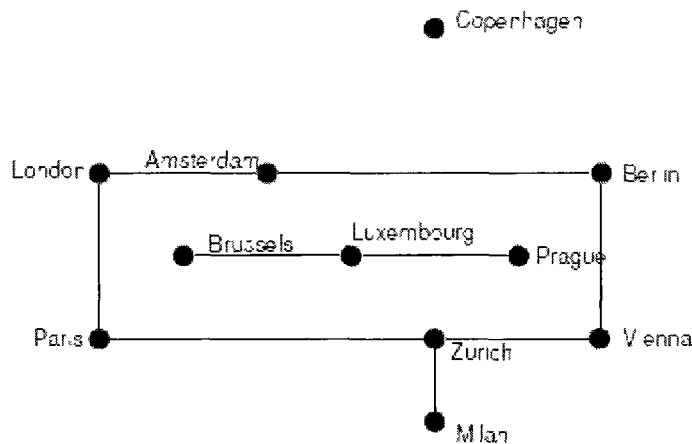


Figure 2.7: Cost239 Network [PAS08]

In the simulation tests, after starting the traffic sources, a link failure is simulated triggering failure detection and switching to alternative routes. To get rid of synchronization effects of hello timers with failure times, the simulations are repeated with different periods of time between the

simulation start and the failure time. The simulation is also repeated for all possible link failures, in order to obtain an average of the effect of different failure locations. In order to characterize the effect of the failure, the sum of the rates of all traffic received at sinks in the network is considered over the time.

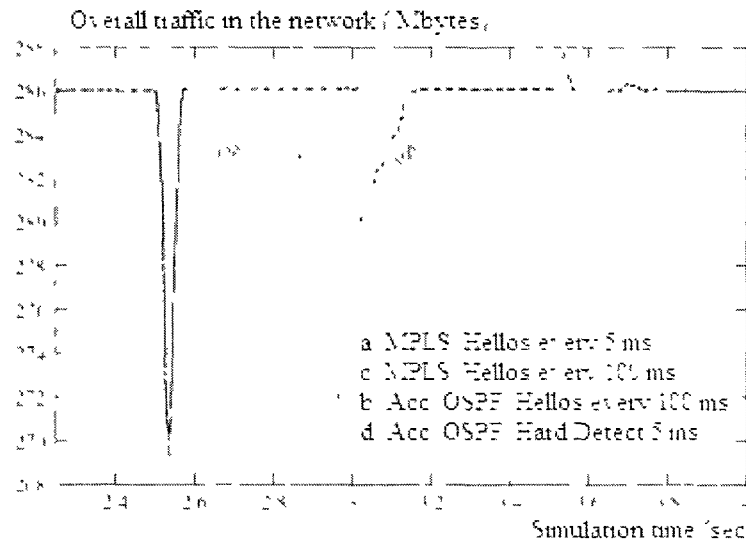


Figure 2.8: Restoration Time [PAS08]

Figure 2.8 shows the affected traffic and the times for restoration for different MPLS protection switching and IP rerouting approaches, both with different timer values for either the RSVP refresh messages or for the OSPF hello protocol. Each curve in the figure shows the sum of all traffic flows in the network.

As a result, Figure 2.8 shows that standard MPLS protection switching a is faster than both OSPF mechanisms. Even, with the same THello and TDead timers as OSPF, MPLS c is remains faster, in the order of 100 ms. These results from the computational effort indicate that the signalling delay from the update of the Forwarding Information Base (FIB) is more time consuming for larger tables of OSPF compared to MPLS.

To conclude, there are two major points to be addressed in order to improve the restoration speed of OSPF re-routing: speed-up of failure detection (hardware failure detection and fast hello protocols), and acceleration of FIB update (the internal router architectures need be improved).

2.8 Experimental Comparison of Fault Notification and LSP Recovery Mechanisms in MPLS Operational testbeds

The authors in [ALB05] compare some recovery strategies in MPLS-TE by conducting experiments in testbeds. The study focuses on alternative notification mechanisms and compares different end-to-end recovery techniques. Moreover, the comparison of different recovery techniques is performed in two different trials, based on commercial routers from Juniper and PC/Linux boxes with a dynamic control plane based on MPLS-TE, as shown in Figure 2.9. Furthermore, two more networks are added. They are Linux PCs acting respectively as traffic generator and traffic collector.

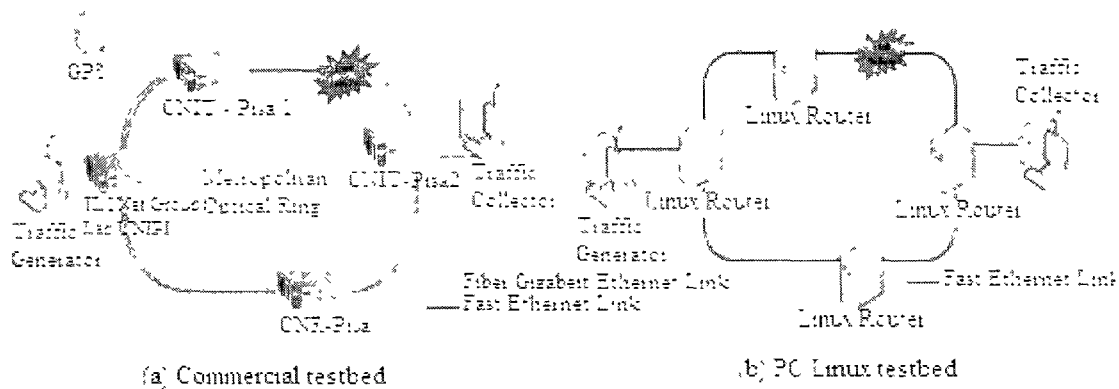


Figure 2.9: Operational MPLS Networks [ALB05]

In both scenarios the traffic generator is directly connected with the ingress router and Constant Bit Rate (CBR) traffic is generated and mapped onto the primary LSP. On the other side traffic is terminated and collected on the traffic collector, which is directly connected with the Egress LER. Traffic traces, with accurate packet time stamping, are collected on both PCs that are synchronized by means of Global Positioning System (GPS) receivers. The experiment begins when a failure is forced to occur on the primary LSP and recovery mechanism starts. The overall recovery time is estimated by the analysis of collected traffic traces.

To compare and distinguish the components such as notification time and recovery operation time, a network analyzer (AdTech AX4000) is used to timestamp packets with a microsecond resolution. Notification is estimated as the time between the delivery of the last packet before the failure and the receiving of the PATH ERR message from downstream router on the Ingress LER.

In order to compute a Mean Recovery Time, the tests are repeated 20 times on the two networks for each mechanism in both network conditions. The results are shown in Table 2.3.

		End-to-End Recovery Mechanisms	Notification Time (ms)	Recovery Operation Time (ms)	Overall Recovery Time (ms)
Commercial testbed	1 LSP	Protection	400,753 (231,428)	27,463 (2,884)	428,216
		Restoration	462,770 (276,290)	31,320 (5,920)	494,090
	10 LSP	Protection	388,610 (240,310)	86,000 (9,930)	474,610
		Restoration	456,730 (264,400)	118,600 (22,270)	575,330
PC/Linux testbed	1 LSP	Protection	39,600 (20,298)	68,700 (15,649)	108,300
		Restoration	40,500 (21,280)	101,300 (28,900)	141,800
	10 LSP	Protection	41,800 (17,570)	113,500 (23,140)	155,300
		Restoration	45,400 (18,330)	138,200 (24,930)	183,600

Table 2.3: Measurements Results [ALB05]

Thus, from the comparison of different recovery schemes, it appears that the proposed flooding-based approach is efficient, fast, and can be implemented within the existing protocols for MPLS traffic engineering, particularly OSPF-TE. Moreover, the results confirm the expectation that the reference of 50 ms of MPLS Fast Rerouting cannot be accomplished by end-to-end mechanisms. On the other hand, the results point towards the possibility of achieving end-to-end MPLS recovery in the order of few hundreds of milliseconds in operational conditions. The objective of using Fast ReRoute is to use point of local repair (PLR) to redirect the effected traffic to the closest working point, using the pre-established LSP. Then the PLR sends an RSVP message to the ingress router to inform it of the failure. In turn, the ingress looks for a new route that does not pass through the affected link.

2.9 Performance Comparison of IP, MPLS and ATM Based Network Cores using OPNET

Asif and Kaosar [ASI06] compare the routing capabilities of three technologies, IP, MPLS, and ATM, based on different performance metrics using OPNET simulator. OPNET is a real time simulator specifically designed for network design and analysis.

For this experiment, the authors use a hypothetical network model as shown in Figure 2.10. The core of the network consists of routers of ATM, IP, and MPLS for each single scenario. All clients (FTP, VOICE and VIDEO) are connected to the leftmost edge router, whereas the single server is connected to the right edge router.

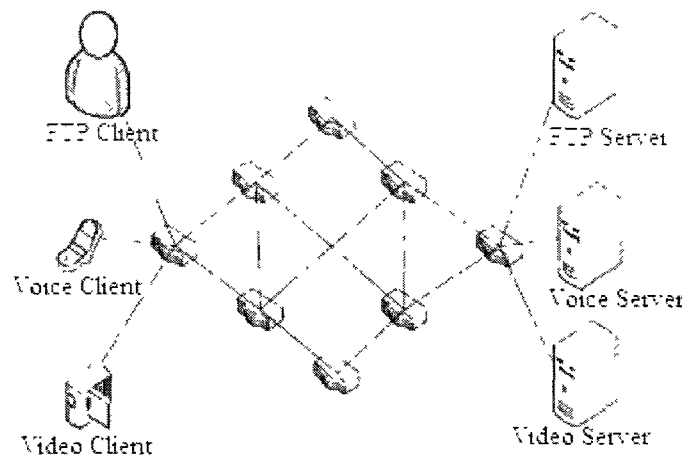


Figure 2.10: Network Topology for Simulation [ASI06]

To compare the performance of the three different technologies, some metrics are considered in the simulations: end-to-end delay, throughput, utilization, and delivery of traffic.

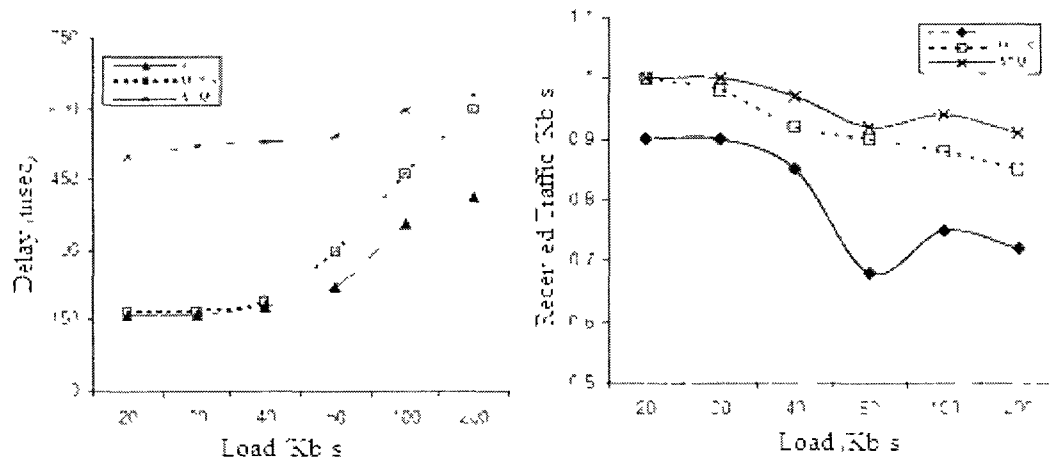


Figure 2.11: End-to-end Delay and Dropped Packets [ASI06]

In Figure 2.11, on the left hand side, the graph shows the delay in IP based network is relatively large compared with ATM, although its behaviour is similar to that of ATM. This is because the absence of virtual connections and MPLS falls between these two. Moreover, the graph on the right hand side shows the connectionless behaviour of IP is responsible for packet drops at the time of congestion or failure. However, it is not always the case that IP routers try to avoid congestion by running congestion avoidance mechanism; although, most of the time, these mechanisms work efficiently. On the other hand, many make full use of connection-oriented nature of the other two technologies, thereby reducing dropped traffic to almost nil.

Even though IP protection mechanism works fine with normal traffic, it is not efficient enough to fulfil the needs of real time applications. In Chapter 4 we provide methods and tools by exploring several scenarios using IP, MPLS, and Fast ReRoute mechanisms, to demonstrate that it is possible to have a competitive convergence time to SONET's convergence time.

2.10 Measurement of Restoration Time from Single-Link Failure in an Optical testbed Environment

The authors [ALA07] present a practical methodology for the measurement of the restoration time from a single-link failure in testbed environment. Two general types of survivability techniques are proposed and implemented: pre-designed protection and dynamic restoration. Pre-designed protection schemes rely on resources (e.g. fibers, wavelengths, switches, etc.) dedicated for protection purposes. The resources are reserved from recovery from failures at either connection setup or network design time, and kept idle when there is no failure. In the case of dynamic restoration, the resources used for recovery are not reserved at the time of connection establishment, but are chosen from the available resources when the failure occurs.

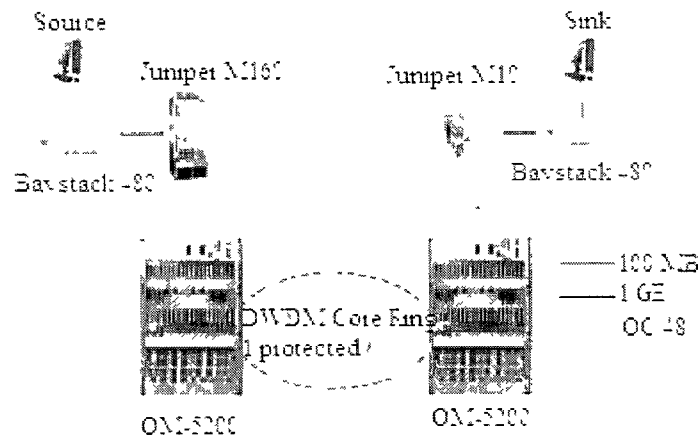


Figure 2.12: Optical Testbed Environment [ALA07]

As shown in Figure 2.12, the testbed includes equipment provided by Nortel Networks and Juniper. First, two Juniper M10 and M160 edge routers and two Nortel Baystack 450 switches are used to represent the Customer Equipment (CE). Second, two OM 5200 switches represent the Metro Dense Wavelength Division Multiplexing (DWDM) network. Finally, two PCs are placed in each of the customer's remote sites. These machines are used as UDP traffic sources/sinks in our tests using Iperf utility.

For the failure scenario, the authors consider the effect of the failure occurring at the primary link of the DWDM ring. In this case, the customer traffic automatically re-routes from the primary to the backup link.

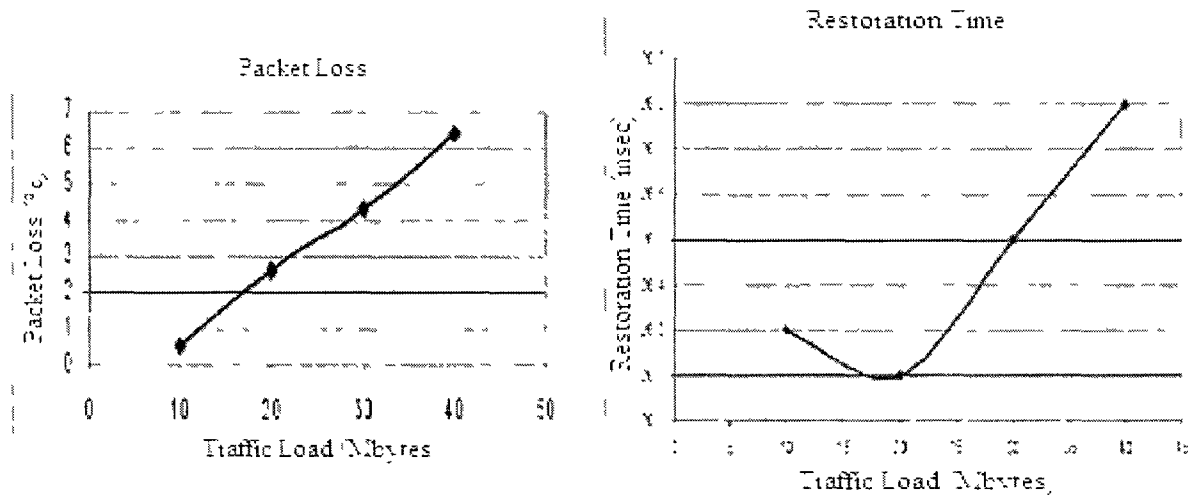


Figure 2.13: Packet Loss and Restoration Time [ALA07]

As shown in Figure 2.13, the packet loss increases as the traffic load increases in the event of a failure, and the maximum restoration time at the load of 40 Mbytes is found to be 26.45 ms. This time is less than the guaranteed restoration time for SONET.

In conclusion, the results from the experiment show it is still possible to achieve faster recovery time without the need for SONET. This is covered in greater detail in Chapter 4.

2.11 Evaluation of IP Fast ReRoute Proposals

In [GJO07] the authors present a comprehensive analysis of some IP Fast ReRoute mechanisms; such are Equal Cost Multiple Paths, Loop Free Alternates, U-Turn Alternates, Tunnels, and Not-via Addresses which are covered in detail in Chapter 3. These mechanisms have been developed by the IETF to activate alternate routing paths which avoid micro loops under node or link failures.

This paper analysis is carried out by evaluating the coverage of those approaches for a variety of inferred and synthetic ISP topologies. In the simulations, the authors test the above mechanisms on a variety of randomly generated topologies, inferred ISP topologies from Rocket fuel project, and real networks such as, ABILENE, GEANT, and BELNET. Their results are based on the

assumption that the topologies are static in terms of link costs which are considered to be symmetric.

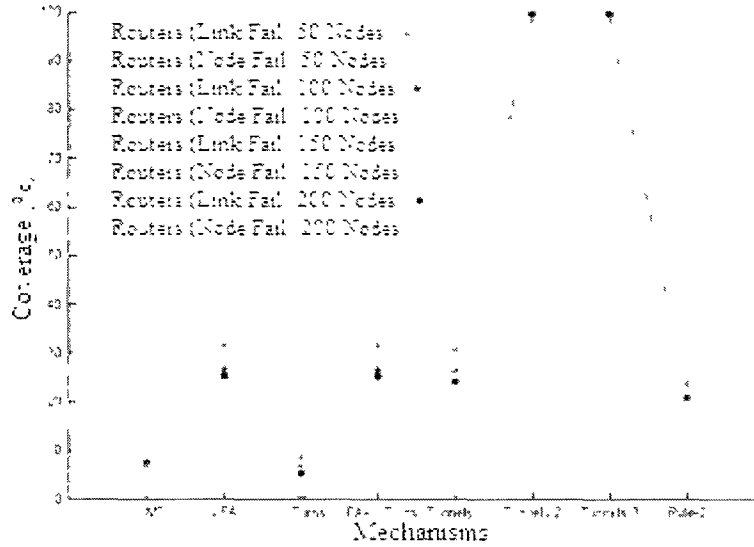


Figure 2.14: Mechanisms Comparison Node/Link Failures [GJO07]

Figure 2.14 shows the coverage achieved for node and link failures for different network size router-level topologies under all mechanisms. It can be noted here that average node degree does not differ by a significant percentage in each of these topologies. The results show that the obtained coverage under node failures is approximately the same for different mechanisms. The same is observed for link failures with the exception of the topology with 50 nodes which has a higher coverage for LFA and U-Turns.

Based on the results, the Not-via- addresses approach has the potential to provide a full coverage to all failures. On the other hand, ECMP, LFA and U-Turns cannot provide full coverage. In any case, ECMP and LFA do not use tunnelling, and therefore, no overhead is incurred. However, U-turn alternates and Not-via- addresses need to use encapsulation, which has an impact on the router performance. In our experiment, we use LFA as a technique for Fast ReRoute. We also study the performance of this technique by making the routing table bigger, which is made by injecting some IGP prefixes. As a result we are able to observe if the convergence time is affected by this increase or stays the same (see details in Chapter 4).

Chapter 3

Network Protection and Restoration Mechanisms

This Chapter explains some of the network level protection mechanisms that have been used in our experiments, namely: layer-3 routing, MPLS global, MPLS-TE FRR, and IP-FRR techniques. Moreover, we discuss how the dynamics of each of the protection techniques in the case of failure, as well as the pros and cons for each.

3.1 Introduction

The ability to detect failure happens in the first step towards providing protection. Physical layer systems provide indication of connectivity loss. One example is packet-over-SONET, which is widely used in network cores and where break in the link is detected within milliseconds at the physical layer. Other transmission media do not have this capability, e.g. Ethernet, which is commonly used in Points of Presence (POPs) [RAH07].

When failure detection is not provided at the physical layer, the mission can be completed at the network layer. This, for example, can be completed using an IGP protocol (OSPF in our case) which sends periodic hello packets to ensure connectivity to the neighbours. When the packets stop arriving, a failure is assumed. However, using this protocol has some disadvantages. Firstly, the architectural limit of IGP hello-based failure detection is 3 seconds for OSPF protocol, and while in most cases, it takes 5 to 40 seconds [ZAB06]. Secondly, handling IGP hellos is relatively complex. Therefore, raising the frequency of generated hello messages creates a burden on the CPU in addition to the higher traffic it generates.

To address these deficiencies, the Bidirectional Forwarding Detection (BFD) [KAT10a] [KAT10b] has been developed jointly by Juniper and Cisco. It has been widely deployed and is

continuing to gain widespread acceptance, as more and more protocols are adapted to use BFD for fast failure detection. BFD is a simple hello protocol designed to conduct rapid failure detection. Its goal is to provide a low-overhead mechanism that can quickly detect faults in the bidirectional path between two forward engines including physical interfaces, sub-interfaces, data link(s), and to a certain extent, the forwarding engines themselves, with potentially very low latency. It operates independently of media, data protocols, and routing protocols.

3.2 Traditional IP Network Recovery (OSPF Layer-3 routing recovery)

In our first scenario we implement a network based on OSPF [MOY98] as routing protocol and established two links between the core routers. To simulate the link failure scenario, we shut down the primary link interface (see Chapter 4 for more details). Once the failure has been detected by the router that is connected to the failed link or adjacent to the failed network node, the adjacent router generates the fault indication signal (FIS) throughout the network in order to inform the rest of the network about the failure. More specifically, the link state protocol propagates the network topology change via OSPF HELLO packets, which is interpreted as an FIS in the case of a failure. Then, upon receiving the notification of the network change, each router in the network quickly re-builds its routing table considering a new path that is bypassing the failure. Then, the network must fully converge to the new path.

Although relying on IGP convergence to overcome the failure has some advantages, such as quick routing around failures and convergence to the remaining topology, it has the downfall that it may take several seconds (5-10 sec) to recover from a link failure [ZAB06]. Source of this delay is due to the Hello protocol that discovers link failures. In order to address this issue, faster failure detection in wired networks has been the focus of several research papers. Such studies attempt to increase network reliability and enable multi-service traffic. For example, [GOY03] suggests reducing the OSPF timer settings to decrease the number of packets lost during link failures in a wired network. However, configuring the IGP to converge quickly can make it very sensitive to minor packet loss, causing false negatives and IGP convergence for no reason. Another proposal is a distributed OSPF scheme [JIA09] for scalable router to schedule routing computation by self-adapting the duration time of network failures. On the other hand, both

[HYU00] and [BAS01] recommend reducing the granularity of link failure timers to milliseconds rather than seconds to improve the stability of wired networks.

As a result, MPLS [ROS01] has been designed to overcome several well-known limitations of traditional IP networks [AWD99a].

3.3 Global Path Recovery (MPLS two paths recovery)

The MPLS global recovery frameworks have been proposed in various researches works [DAS02] [SAN01] [SUR01] [CHA02], which are achieved by using Label Switched Paths (LSPs) in MPLS core networks. There are two types of paths: a primary path, used to carry the traffic in normal operation; and a pre-established secondary path or backup path used to take over when failure occurs.

Figure 3.1 shows one dashed primary LSP along the path from the ingress node to the egress. A dotted backup LSP takes another path and is ready take over. Although it does not carry any data, it might carry control traffic. When an error in the primary LSP occurs, it is reported back to the ingress LER using notify message (RSVP messages) [BRA97]. As a response, the data is immediately switched to the backup LSP, maintaining traffic flow [SHA03].

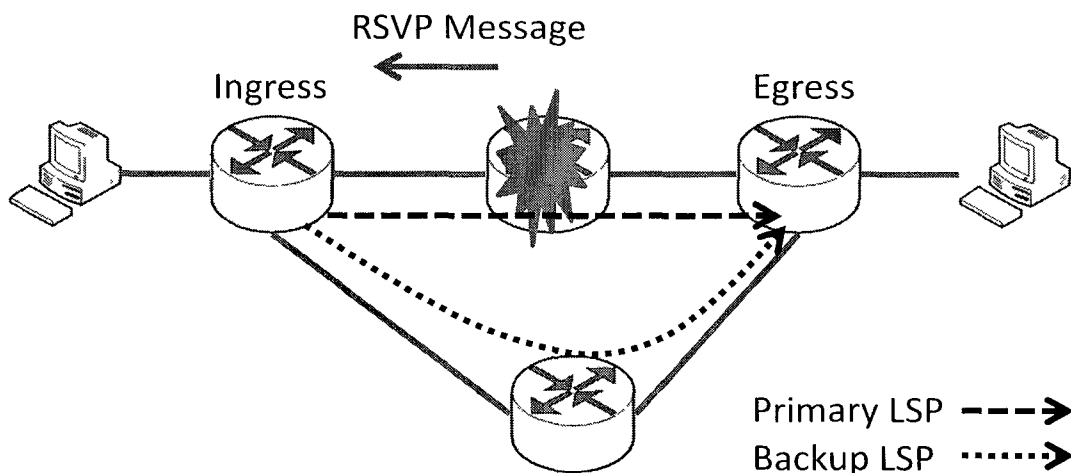


Figure 3.1: MPLS two Paths Recovery

Using this recovery scheme is faster than using IGP [RAH08]. This has the advantage that network resources are not tied up by the backup LSP. However, because path protection is provided by the ingress router for the entire path, it increases the failure over time and it is subject to the prospect of no resources being available when the backup is needed. (This varies from simply propagating the error from the point of detection to the point of repair, to the full signalling of the backup LSP. To solve this, a local protection or MPLS-TE FRR can be used to go over this delay and switch over quickly.

3.4 Local Path Protection (MPLS-TE FRR)

MPLS-TE [NAD10] [AWD99b] provides a mechanism to create a traffic-engineered connection-oriented path, called Label Switched Paths (LSPs), between IP routers. Before a path can be calculated, a view of network resources is determined using layer 3 (Network Layer) routing protocols, such as OSPF or IS-IS, with traffic engineering extensions. Based on this network view, a path is selected for an LSP using Constrain Shortest Path First (CSPF) algorithm. Once the path has been calculated, the path is signalled using RSVP-TE [PAN05] based on the extension of the reservation protocol to support traffic engineering. This is a soft state signalling protocol, requiring periodic refresh messages to ensure links are still active. The CR-LDP is an extension based on Label Distribution Protocol (LDP). CR-LDP [JAM02] also allows for reservation of bandwidth and is a hard state signalling protocol, i.e. it does not require refresh messages to maintain an active link [FOO03].

The goal of Fast ReRoute protection is to minimize the time during the failure. Thus, it is logical to apply protection as close as possible to the point of failure. The idea of Fast ReRoute (a local protection recovery scheme) is simple. Instead of providing protection at the head end (Ingress) for the entire path, the traffic around the point of failure is rerouted.

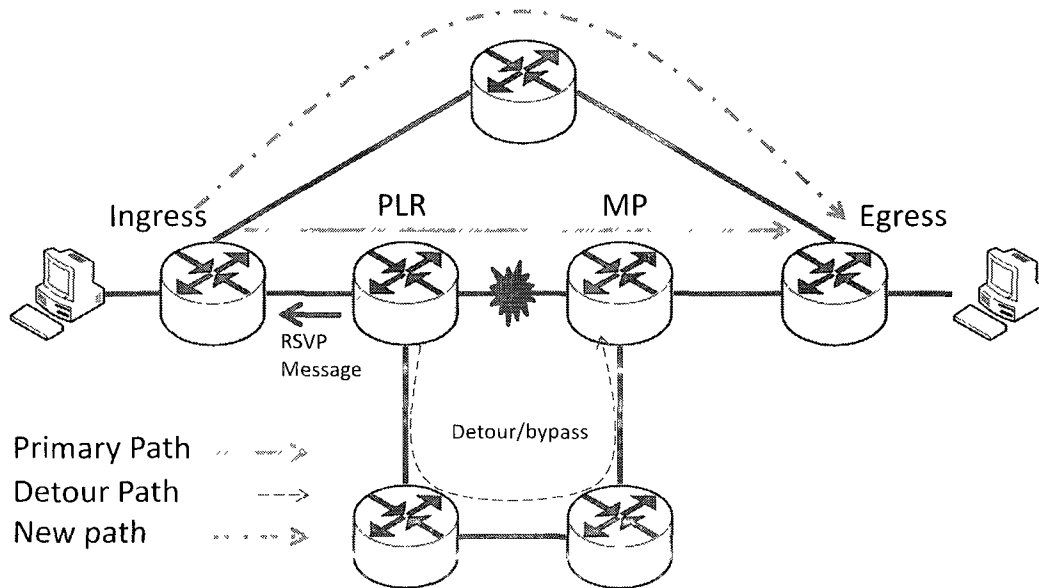


Figure 3.2: MPLS-TE FRR Protection

As shown in Figure 3.2, the LSPs affected by a failure are locally rerouted by the node immediately upstream to the failure. This node is called the point of local repair (PLR) and the node where the backup tunnel terminates is called the merge point (MP). As a result, the traffic is quickly rerouted around the point of failure. Thus, this technique is called Fast ReRoute. The objective of this technique is not to keep the traffic on the detour until the link recovers, but rather to keep it long enough for the LSP head end to move the LSP to a new path that does not use the failed link as shown in Figure 3.2. This displays the very significant advantage of providing a faster recovery time in case of link or node failure. However, this protocol uses RSVP as a signalling protocol, which increases the traffic load of the network [LAK05] [HUA07a].

An alternative Fast ReRoute scheme that provides a local recovery against link or node failure is IP Fast ReRoute which employs the backup routes in pure IP networks. This scheme is discussed in Section 3.5.

3.4.1 LDP over RSVP

In our MPLS-TE scenario, we use LDP over RSVP. This is because the MPLS-TE technology consists of building two, multiple, or even full-connected LSPs in the network to schedule the network traffic, thus achieving the balance of traffic in the network. As a result, the network traffic is balanced by transferring traffic from heavily loaded links to the links with less traffic. The MPLS-TE has a sharp edge in addressing the network congestion.

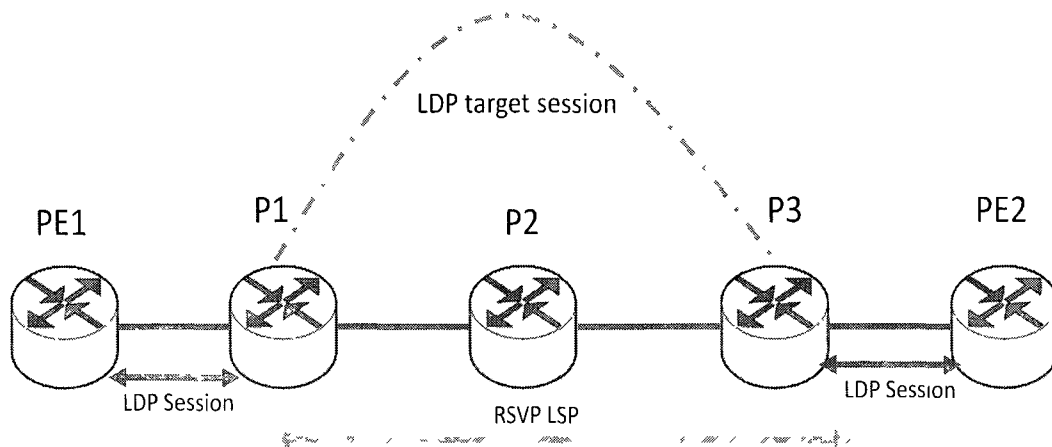


Figure 3.3: LDP over RSVP

Currently, it seems very difficult to deploy the TE throughout the entire carrier network. Therefore, any Internet Service Provider (ISP) can plan a core RSVP area and deploy TE in this area (P1 to P3) while enabling LDP out of the area (PE1 to P1) and (P3 to PE2) as shown in Figure 3.3 [HUA07b].

3.5 IP Fast ReRoute Protection

3.5.1 Introduction

A number of years ago, Shand and Bryant of Cisco Systems [SHA10a] presented an Internet-draft to the IETF called IP Fast ReRoute Framework. In January 2010, it became an RFC that proposes an FRR solution without MPLS-TE.

This approach computes backup routes that allow the failure to be repaired locally by the routers detecting the failure. This is conducted without the need to inform other routers of the failed network, and its reroute paths are intended only as short-term detours around a failure while the network's IGP calculates a new route. As a result, the recovery time can be reduced to a shorter time to detect the adjacent failure and invoke the backup routes providing less than faster recovery time.

In IP-FRR [SHA10a], when a failure occurs, it is detected in three ways: physical detection, e.g. loss of light; protocol detection, e.g. Bidirectional Failure Detection protocol (BDF); and routing protocol detection, e.g. use of hello packets. Following detection, the traffic is transmitted through one or more repair paths.

3.5.2 Recovery Mechanisms in IP-FRR

3.5.2.1 IP-FRR using Loop Free Alternate

The basic idea behind Loop Free Alternates [ATL08][RAJ07] is to use a pre-computed alternate next hop in the event of a link failure to allow traffic to be routed through this alternate next hop when a failure is detected. Thus, the traffic flow is not disrupted and the network can converge in the background. Once the network has converged and the routing tables are updated, the traffic then flows along the newly calculated primary next hop.

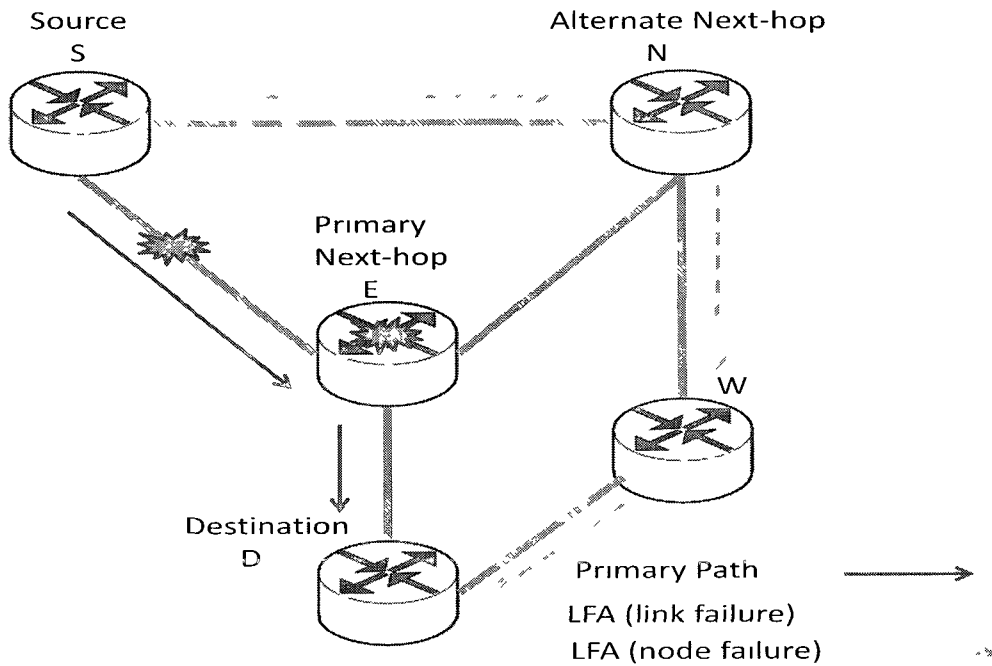


Figure 3.4: Loop Free Alternate Scheme

Figure 3.4 shows the alternative paths from node (S) to the destination (D) via node (N) which guarantees the traffic does not traverse the failure, i.e. the failed link or node is not included in the alternative path. IP Fast ReRoute specifies a condition for Link-protecting alternates and a more restrictive condition for Node-protecting alternates. In order guarantee loop-free alternates, these conditions must be maintained as follows:

$\text{Cost}(N, D) < \text{cost}(N, S) + \text{cost}(S, D)$ for Link-protecting alternates

$\text{Cost}(N, D) < \text{cost}(N, E) + \text{cost}(E, D)$ for Node-protecting alternates

3.5.2.1.1 Recovery Process

In IP-FRR LFA scheme the source node computes the shortest path to the destination node, and then the source node chooses one of its neighbour nodes as a primary next-node to the destination. Without IP Fast ReRoute, the source node chooses only one next-node to be in its path to the destination. With IP Fast ReRoute, the source node also looks for an alternate next-node to be in the secondary path to the destination.

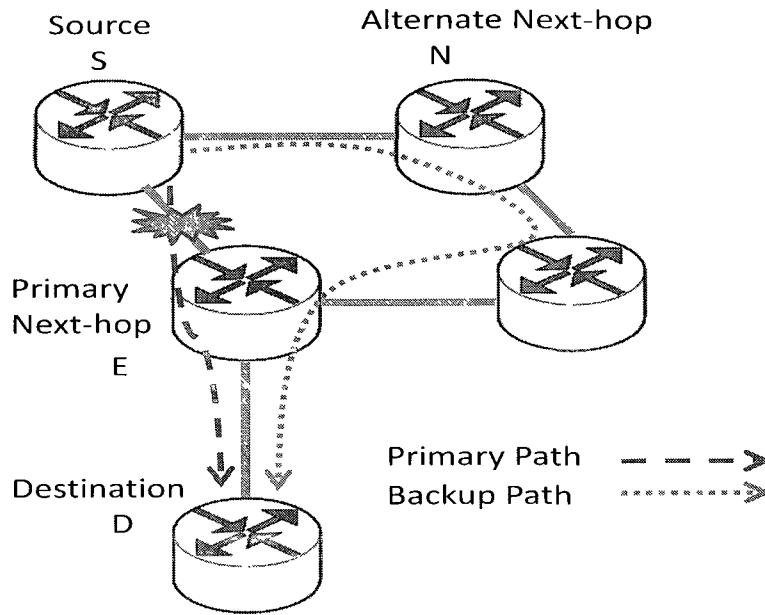


Figure 3.5: Loop Free Alternate Link Failure

Figure 3.5 shows the recovery process in case of link failure. When the link between the source node (S) and the primary next-node (E) fails, both the source node (S) and the primary next-node (E) detect it. On detecting the failure, the source node stops sending traffic to the destination node (D) via the failed link (primary next-node), and instead, it sends the traffic to the source's pre-computed alternate next-node (N), which is the backup link, until a new short path first algorithm (SPF) is run and its results are installed. However, this creates a forwarding loop because the shortest path from the source to the destination is via the primary next-node.

Therefore, to avoid such routing loops, the following Loop-Free Condition (LFC) should be satisfied: $\text{Cost}(N, D) < \text{Cost}(N, S) + \text{Cost}(S, D)$. LFAs are pre-computed and installed in the Forwarding Information Base (FIB) of a router for each destination.

3.5.2.2 IP-FRR using Not-via-addresses

Not-via-address [SHA10b] [RAJ07] is an additional address assigned to each interface in the network. The objective of this scheme is to protect the failure of a next-node (P) or of its adjacent links by moving the affected traffic around node (P) or the failed node to the next-next hope (NNHOP) (B) using Not-via address. Moreover, the backup path must not contain the failed node (P) which is not the case with normal IP forwarding because node (P) is on the shortest path from the source node (S) to the destination node (D). Therefore, special “Not-via-addresses” Bp are introduced such that packets addressed to (Bp) are forwarded to (B) not-via node (P) as shown in Figure 3.7.

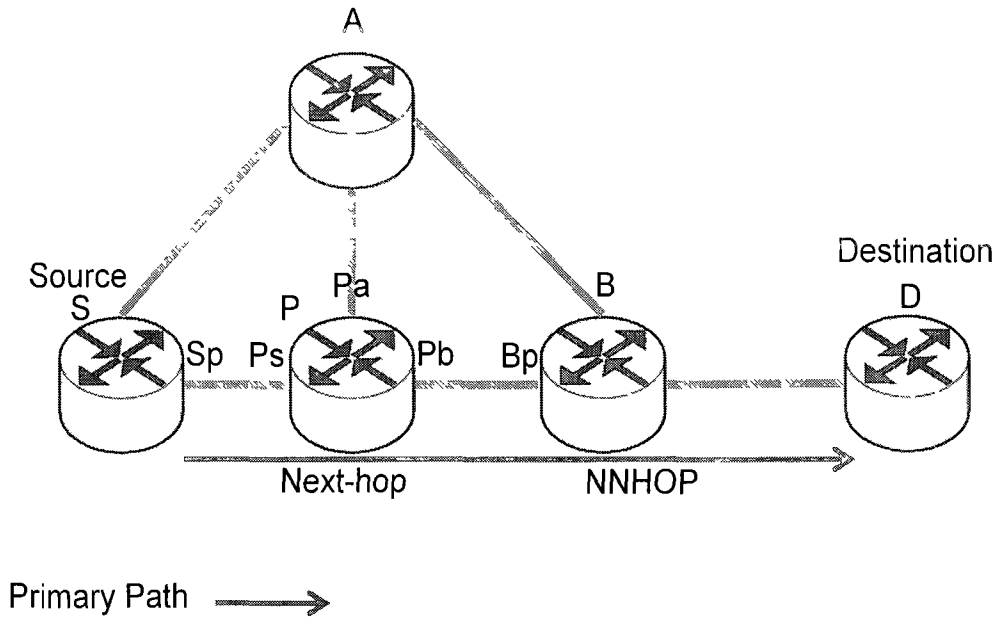


Figure 3.7: Not-via-addresses Scheme

As can be seen, there is a huge number of Not-via-addresses imposing a considerable load on IP routers. However, the author of [ENY09] provided a solution (lightweight Not-via) that substantially decreases the number of additional addresses. Moreover, they improve the complexity of computing the detours to strict linear time from the worse than quadratic complexity of Not-via. Also, a combined application of Not-via-addresses and LFA can achieve 100% single failure coverage, while using simple LFAs wherever possible [MEN10].

3.5.2.2.1 Recovery Process

As shown in Figure 3.8 below, in normal operation, the source node (S) sends packets to the destination node (D) via (P). When failure occurs such as the next hop (NHOP) (P) fails or the link from (S) to (P) fails. Then the node (S) encapsulates its packets that go to (P) in another IP packet addressed to the NNHOP, which in this case is (B), using the Not-via- address (Bp). These packets are forwarded from node (S) over node (A) to node (B), which is the shortest path around node (P). NNHOP, which is (B), performs decapsulation and forwards the repaired packets towards the destination node (D) [SHA10b] [RAJ07].

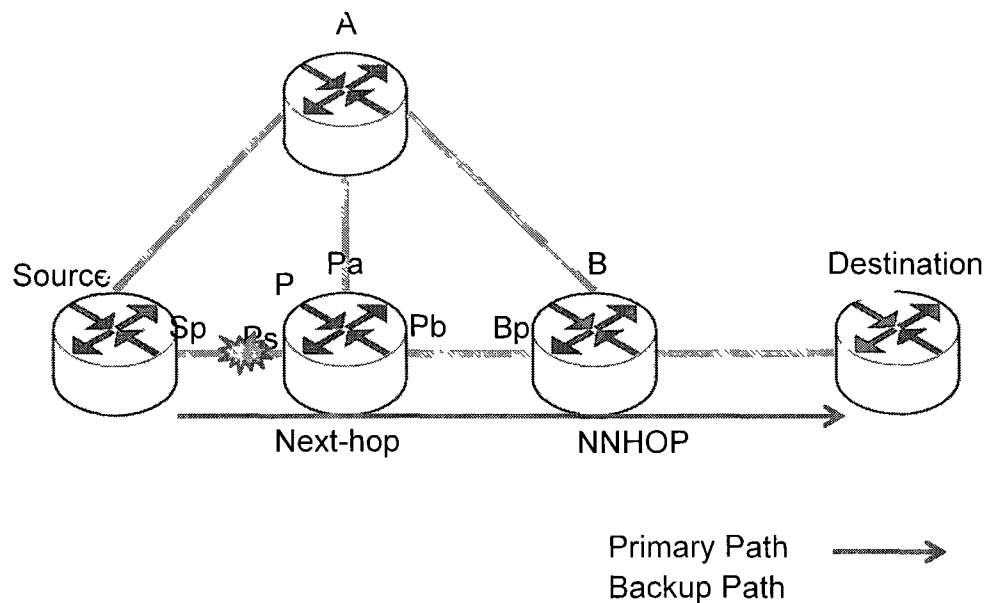


Figure 3.8: Not-via-addresses Link Failure

3.5.2.2.2 Pros and Cons

1. It guarantees 100% failure coverage for single node and link failures [GJO07].
2. Encapsulation and decapsulation are performed on the source router and the NNHOP router which may have a performance impact on router hardware.
3. It requires additional entries in the forwarding tables which increases routing table size.
4. It deviates the traffic around the failure back onto the original path which makes it longer than LFA scheme.

3.5.2.3 IP-FRR using U-turn Alternate

U-turn alternate [SHA10a] [RAJ07] allows packets to travel back on the same link two times. This procedure needs a special mechanism that allows this U-turn forwarding. The topology in Figure 3.9 is an example of a U-turn alternate scheme.

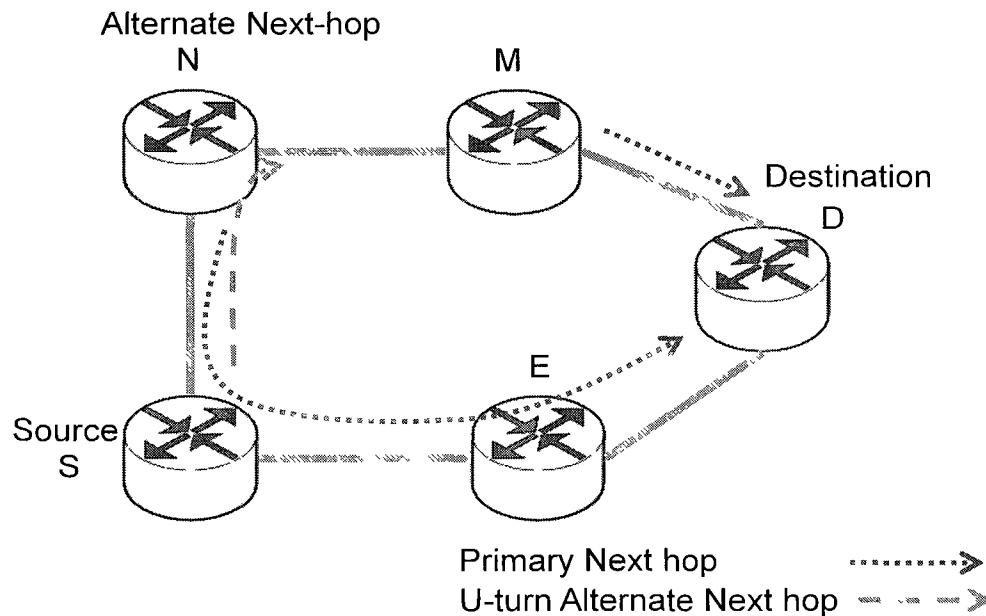


Figure 3.9: U-turn Alternate Scheme

In Figure 3.9, there is no loop-free alternate for node (S) to use to reach node (D) because the costs are such that node (N) uses node (S) as its primary neighbour; therefore, if (S) wants to send the traffic to (N), it would loop back to S. If both (S) and (N) support the U-turn mechanism, then (S) could use (N) as a U-turn alternate. Traffic destined to node (D) that is sent by node (S) to node (N) would be forwarded by node (N) to node (M), N's loop-free node-protecting alternate.

Such alternates may be used when no other loop-free alternates exist, and when the following conditions are met:

1. N must have S as its primary next hop towards D.
2. N must have a node-protecting alternate M for the destination D with respect to node S.
3. If covering node fails, the path from M to D must not contain the failed node E.

Combining U-turn Alternate scheme and Loop-free Alternates can be sufficient to protect up to 90% of the link or node in case of failure. Furthermore, adding protection tunnels to these two basic techniques is sufficient to achieve a full coverage [FRA05b].

3.5.2.3.1 Recovery Process

If an alternate next-hop is available, the router should redirect traffic to the alternate next-hop when the primary interface has failed. In order for this to occur, detection of that failure has to occur within an appropriately short period.

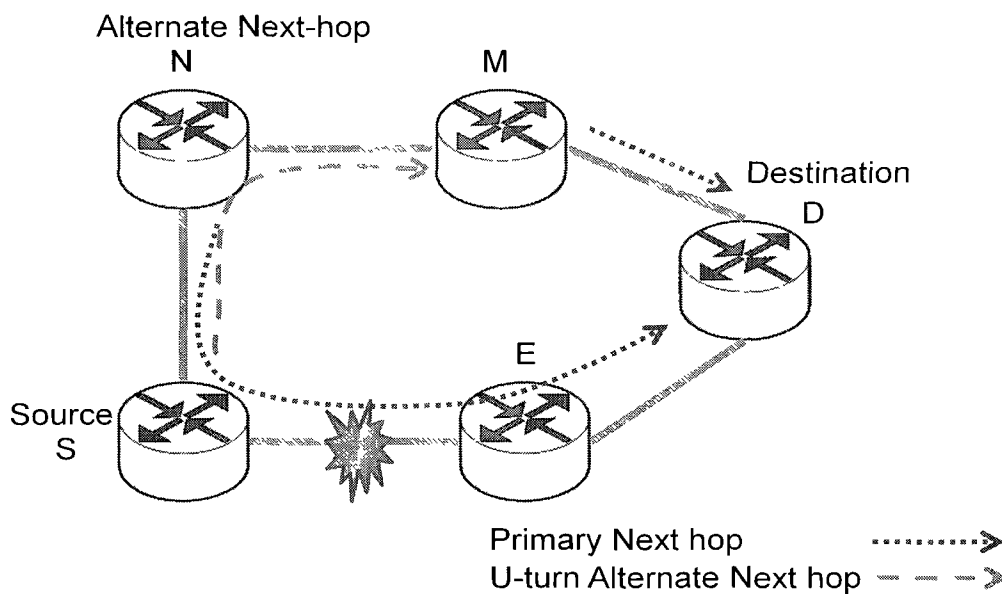


Figure 3.10: U-turn Alternate Link Failure

In Figure 3.10, node (N) provides a U-turn alternate to (S) via the loop-free node-protecting alternate (M). If node (S) wants to send the traffic to node (N) after the link between node (S) and node (E) fails, it would loop back to (S). If both (S) and (N) support the U-turn mechanism, then (S) could use (N) as a U-turn alternate. Thus, (S) could send traffic to (N), and (N) would forward it to (M), (M) would then forward it to (D). Here, M is N's loop-free node-protecting alternate.

3.5.2.3.2 Pros and Cons

1. The U-turn alternate procedure considers only single hop U-turns. When the U-turn is performed, the receiving node must have an LFA path to the destination. Otherwise the neighbour cannot be used as a U-turn neighbour. Therefore, the protection coverage is still topology dependent.
2. Packets sent to the alternate next-hop require adding a single U-turn label. This needs to be removed at the next hop.
3. This mechanism does not work in the ring topologies and only works for the networks such as triangles of nodes. This is because it only provides limited protection coverage.

3.5.2.4 IP-FRR using Tunnels

This mechanism [BRY07] uses IP tunnels for achieving limited directed forwarding. The directed forwarding refers to the ability at the tunnel end point, which is the decapsulation point, to specify which next-hop should be used for forwarding the affected packets, rather than the normal one indicated by a lookup of the destination address in the FIB. Therefore, directed forwarding permits a router to specify the tunnel egress where traffic should be forwarded. This tunnel encapsulation could be an MPLS label encapsulation. When IP or MPLS tunnels are used as repair paths, the repair strategies operate on the basis that if a packet can somehow be sent to the other side of the failure, or can be dropped at a point where it can reach the other side, it subsequently proceeds towards its destination exactly as if it had traversed the failed component.

3.5.2.4.1 Recovery Process

As shown in Figure 3.11, when a node (S) detects an adjacent failure it uses a pre-computed set of repair paths to bypass the failed link between node (S) and node (E). A tunnel is used to carry traffic to a router, called the tunnel endpoint, where loop free alternate paths to the destination router exist, using normal forwarding. All packets are encapsulated by node (S) and routed towards the tunnel endpoint which is node (E) in this case. The tunnel endpoint then decapsulates the packets and forwards them, according to normal routing, towards the destination D as shown in Figure 4.9. Note that the tunnel does not have to go from S to E; it can terminate at any router in the network, provided that S can be sure that the packet proceeds correctly to its destination

from that router. This scheme distinguishes between a link failure and the failure of its neighbour (node failure). Repair paths are calculated in anticipation of adjacent router failure.

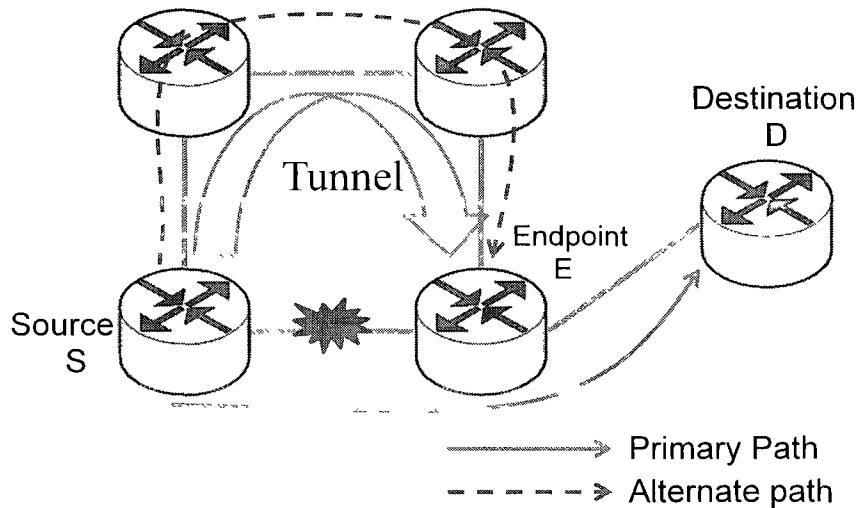


Figure 3.11: Using Tunnels Link Failure

3.5.2.4.2 Pros and Cons

1. A tunnel-based backup procedure requires the following enhancements to routing protocols:
 - a. The ability to advertise IP-FRR capability.
 - b. The ability to advertise tunnel endpoint capability.
 - c. The ability to advertise directed forwarding identifiers.
2. Multiple concurrent failures other than those that occur due to the failure of a single router are not addressed.
3. Shared risk group protection is not considered. However, this procedure can trivially be trimmed to exclude SRLG traversals, but the resultant connectivity may require a secondary repair analysis.
4. It relies on complete topological information from the link state routing protocol; they only work within a single link state flooding domain.

3.6 Summary

Building a network that can survive a failure is not simple. The difficulty lies in the fact that the network is a complex architecture of network elements, such as nodes and links, which can fail at any time. As discussed earlier, many networks, such as IP, and MPLS, are originally designed with little attention paid to rapid recovery from failures.

Protecting any network element against failure can be accomplished by using the inherent properties of IP and IP routing. The techniques of protection switching derived from experience in other signalling protocols, or Fast ReRoute methods are developed specifically for MPLS. Each has its own specific characteristics, advantages and disadvantages. Consequently, the method used must be chosen in accordance to the requirements of the network and the users.

If very rapid repair is needed, for example for real-time traffic such as voice and video, then Fast ReRoute schemes most likely provide the best solution. If quick repair with the possibility of sharing backup resources is desired, then MPLS two path recovery technique can be chosen. If repair time is not crucial and network resources are limited, then an IGP protocol alone may be sufficient.

Chapter 4

Experiments and Analysis

In this chapter, we present the methods and tools that we have employed in carrying out our experiments. The experiments consider the performance of two FRR techniques in IP and MPLS networks. All experiments have been performed in the research testbed of the Optical Network Research Lab (ONRL) at the University of Ottawa. The routers that are used in our tests network consist of real and, emulated ones. In first three scenarios, we successively evaluate the performance of the IP and MPLS recovery mechanisms such as, IGP recovery, MPLS recovery, and MPLS local protection. In the second part, we also evaluate the performance of Fast ReRoute mechanisms in cases of IP and MPLS core networks, such as MPLS-TE FRR and IP-FRR.

In Section 4.1, we discuss in more details the hardware used to conduct our experiments. In Section 4.2, we assess the performance of three different recovery schemes by simulating link failures and measuring packet loss and convergence time for each case. The convergence time in these scenarios is considered as a recovery time or end-to-end delay. In Section 4.3, we evaluate the performance and analyze the results of tests for the IP-FRR and MPLS-TE FRR mechanisms. We consider the effect of increasing a number of parameters such as the number of LSPs affected by a single failure, and the number of routing prefixes on the overall convergence time and traffic loss. The convergence time in this case is considered as the summation of the recovery time and the time taking to signal a new path.

4.1 Hardware used for the Experiments

The hardware used in our tests consisted of Nortel and Juniper-based routers and switches, as well as Linux-based emulation routers. In addition, two IBM PCs that are equipped with a Gigabit Ethernet (GE) ports are placed in each of the customer's remote sites. The clocks on both

PCs are synchronized using Windows Network Time Protocol (NTP) server. These PCs are used as traffic sources/sinks in our tests (see Figure 4.1)



Figure 4.1: Testbed used for our Experiments

④ Juniper Routers

- Both M10 and M160 routers are running JUNOS version 10.2 Internet Software supporting MPLS, OSPF, BGPv4 and VPNs implementation
- Routers are connected using Packet Over Sonet (POS) OC-48 and GE links with capacity of 2.45 Gbps and 1 Gbps, respectively
- Each router is uniquely assigned a loopback interface address, as well every interface of each router is assigned a unique IPv4 address from the private address space

- **Nortel Optera Metro 8000s (OM-8000s)**

- Both Optera Metros are running Software release OM 2.0.1.010.2 that supports P2P layer 2 VPNs over MPLS, LDP, RSVP-TE, and VPLS.
- Both are equipped with 4 x GE ports.
- They are designed to provide high-density connectivity for mid-size and large enterprise represents the provider's side.
- Each router is uniquely assigned a loopback interface address, as well every interface of each router is assigned a unique IPv4 address from the private address space.

- **Nortel BayStack Gigabit Ethernet Switches**

- Both BayStack switches are running BayStack Operating system Switching Software (BoSS) v.3 which offers highest level of security with featuring including Secure Shell SSH and Simple Network Management Protocol (SNMPv3).
- 24 x Network – Ethernet 100Base-TX - RJ-45.
- They include two built-in GE ports.
- They are designed to provide high-density desktop connectivity for mid-size and large enterprise customers wiring closets.
- The Baystack switches are connected to the MPLS domain using 100Base-T Ethernet connections.

- **IBM PCs Host Computers**

- Intel Core 2 Duo CPU E8200, 2.66 GHz processors, 1.95 GB RAM.
- Windows XP SP2-based and is equipped with Gigabit Ethernet (GE) ports.
- They are used to run the Iperf application tool.

4.2 Network Recovery Scenarios

4.2.1 Description of Testbed Hardware

Figure 4.2 shows the testbed that has been used to carry out the experiments. It consists of several Nortel and Juniper-based equipment. The testbed is configured as follows:

First, the M10 and M160 routers are equipped with GE and POS OC-48 interfaces that are enabled for IP/MPLS to handle and redirect multiple traffic flows through two paths to designated route in case of failure.

Second, the two Nortel layer-2 switches (BayStack 470-24T) are equipped with 2x GE uplink ports that are used to connect the PC end computers in our testbed.

Finally, the PCs are equipped GE ports and placed in each of the customer's remote sites. The clocks on both PCs are synchronized using Windows NTP server and are used as traffic sources/sinks in our tests.

Furthermore, the Iperf traffic generator tool is used to generate multiple UDP traffic flows to measure recovery time and packet loss. The Iperf tool is capable of measuring a number of parameters, including bandwidth, end-to-end delay, and jitter.

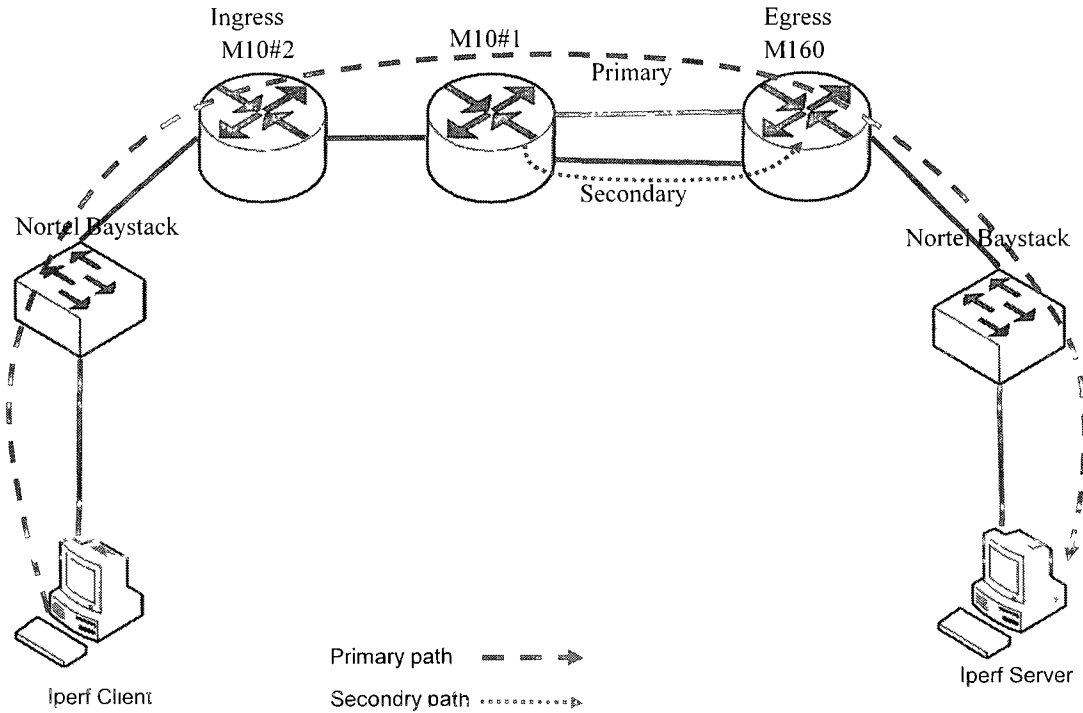


Figure 4.2: Network Topology used for Network Recovery Scenarios

4.2.2 Experiment Methodology

In our experiments, the Iperf generate traffic that is used to measure the packet traffic loss resulting from the switchover of traffic from primary to protected paths following a link failure event in the network. The M160 and M10 routers are configured as an IP based network relaying on OSPF protocol as illustrated in the first scenario, MPLS based network with no FRR in the second scenario, and MPLS-TE FRR in the third scenarios.

The customer traffic (Iperf UDP flows) are carried through the provider's core (IP or MPLS) network. At the provider edge core network, traffic is aggregated to pass along over POS OC-48 link acting as the primary link, while another GE link is provisioned as backup.

In our tests, we focus on performance metrics such as packet loss and recovery time of the network following a link failure in the network when implementing any of the IP and MPLS recovery mechanisms. In the first case, a pure IP core network is built by enabling Open Shortest Path First (OSPF) routing protocol on all the core and non-core routers. In the second case, an MPLS-based network is built using LSR routers and traffic is transported over unprotected Label

Switched Paths (LSPs) from the source to the destination. In the third scenario, MPLS-TE FRR functionality is deployed as a recovery mechanism to protect traffic in cases of failure.

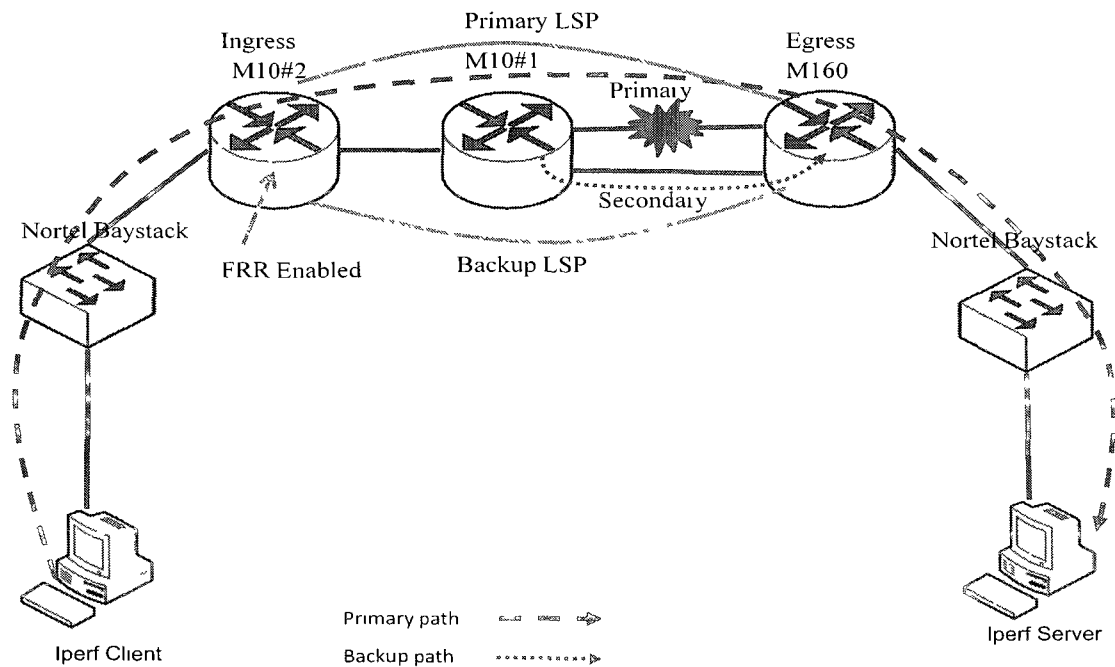


Figure 4.3: Primary Link Failure in Network Protection Scenarios

Figure 4.3 shows the two Baystack 470 switches which configured as layer-2 switches carrying untagged traffic from PCs running the Iperf tool. For this test, two paths are provisioned acting as primary, through the POS OC-48 link, and the other as a backup, through GE link. The Iperf tool is configured to run on both PCs in order to measure packet loss and convergence time for traffic sent across the provider core network.

4.2.3 Three Network Recovery Scenarios

This section describes three different scenarios that are implemented in the research testbed. For each test, the Iperf tool is configured to send traffic at rate of 10kbps for the duration of 20 seconds.

4.2.3.1 Scenario#1 (Layer-3 OSPF routing recovery)

In this scenario, we enable OSPF as an IGP routing protocol on all interfaces of each router in the network. A primary POS OC-48 link is provisioned between M10 and M160 routers. Another parallel link between M10 and M160 routers over GE is also provisioned to be used in the case of failure of the primary link as shown in Figure 4.3. The OSPF weights on the links are biased such that routing picks the POS primary link before the link failure.

To simulate the link failure, the POS primary link is failed by administratively taking it out-of-service while the backup link is still connected. We repeat this test five times and get the average. It is observed that around 9 seconds is needed for OSPF routing to converge and reroute traffic over the backup link. In this case, almost 98.87 ms of traffic loss is observed during convergence before traffic resumed as demonstrated in Table 4.1.

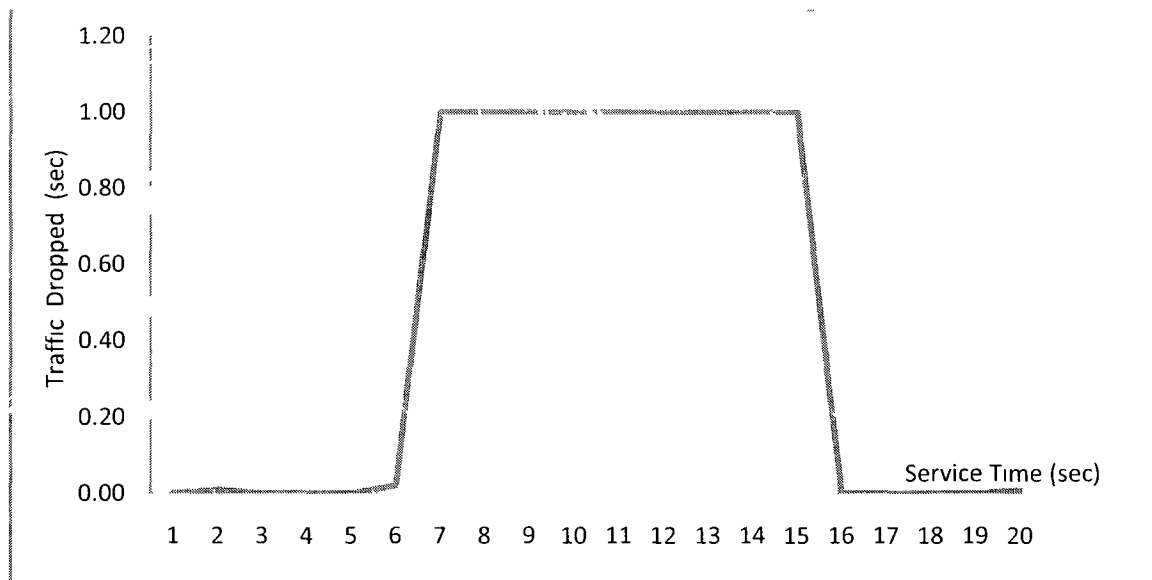


Figure 4.4: Traffic Dropped in Layer-3 OSPF routing Recovery Scenario

Figure 4.4 shows the traffic disruption in real-time in the interval [6, 15] after the link is failed and before routing convergence completes. For most customer traffic or application, such IGP recovery process is very slow and unacceptable.

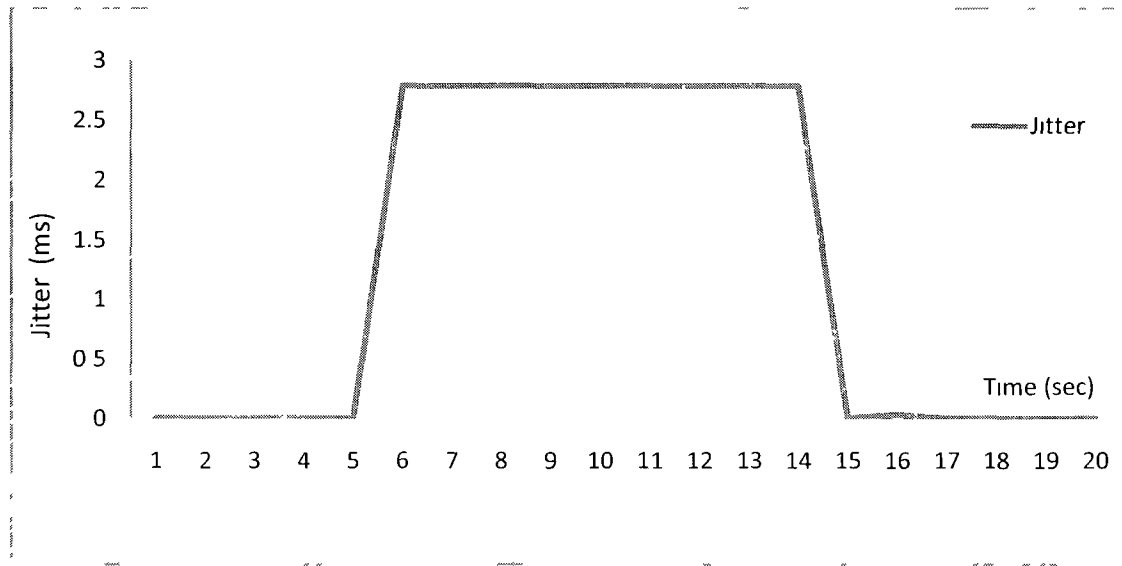


Figure 4.5: Jitter in Layer-3 OSPF routing Recovery Scenario

Figure 4.5 shows the delay jitter, as measured using Iperf tool, increases to almost 3 ms when the failure is injected, while no jitter is noticed before the link is failed or after convergence completes. During convergence some of the traffic gets rerouted onto sub-optimal paths while others onto lower delay links. This results in variation in the delay experience by consecutive packets of the flow producing the observed jitter.

4.2.3.2 Scenario#2 (MPLS two Paths Recovery)

In this scenario, we have enabled the provider's core network to support MPLS. Also, we have create bidirectional dynamic Label Switched Paths LSPs between the M10 and M160 LSR routers that traverse over the primary POS OC-48 link, and another set of explicitly defined LSPs that traverse over the GE link. The two bidirectional tunnels of capacities 500 Mbps are created on each of the LSPs, and act as the primary and standby backup as shown in Figure 4.3.

To simulate the link failure, the POS primary link is failed by administratively taking it out-of-service while the backup link is still connected.

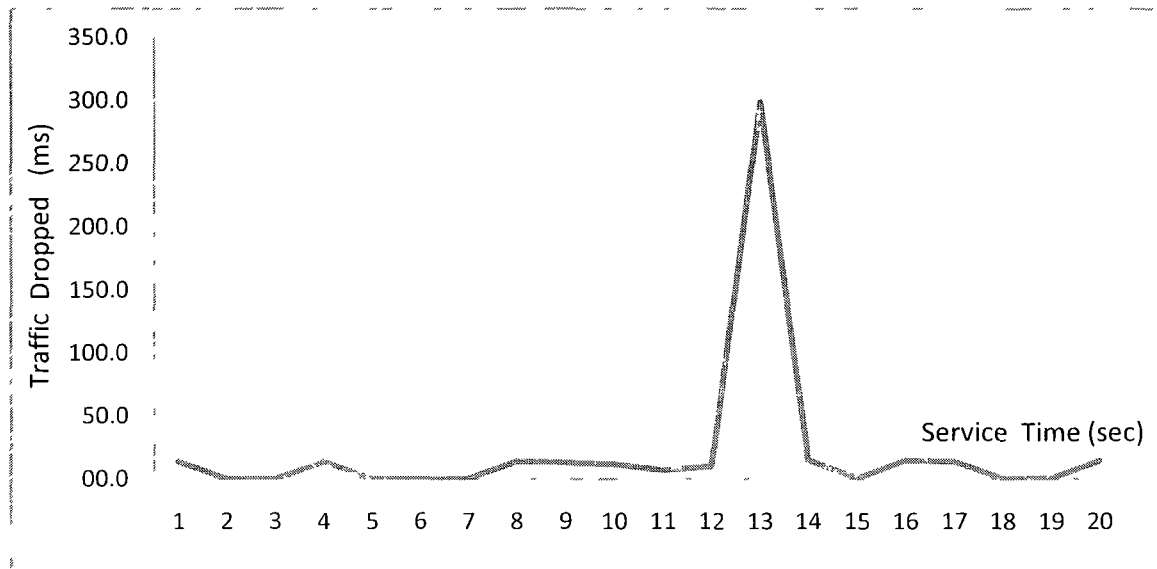


Figure 4.6: Traffic Dropped in two Paths Recovery Scenario

As shown in Figure 4.6, the packet loss decreases from the first scenario to reach 298 ms traffic loss between seconds 12 to 14, resulting in a 10.3 ms as calculated in Table 4.1 total packet or traffic loss delay. This indicates using RSVP signalling protocol gives the MPLS core network the ability to recover faster than using IGP routing protocol. However, the PLR still sends the RSVP Error messages all the way back to the ingress router, which adds more delay to the recovery process.

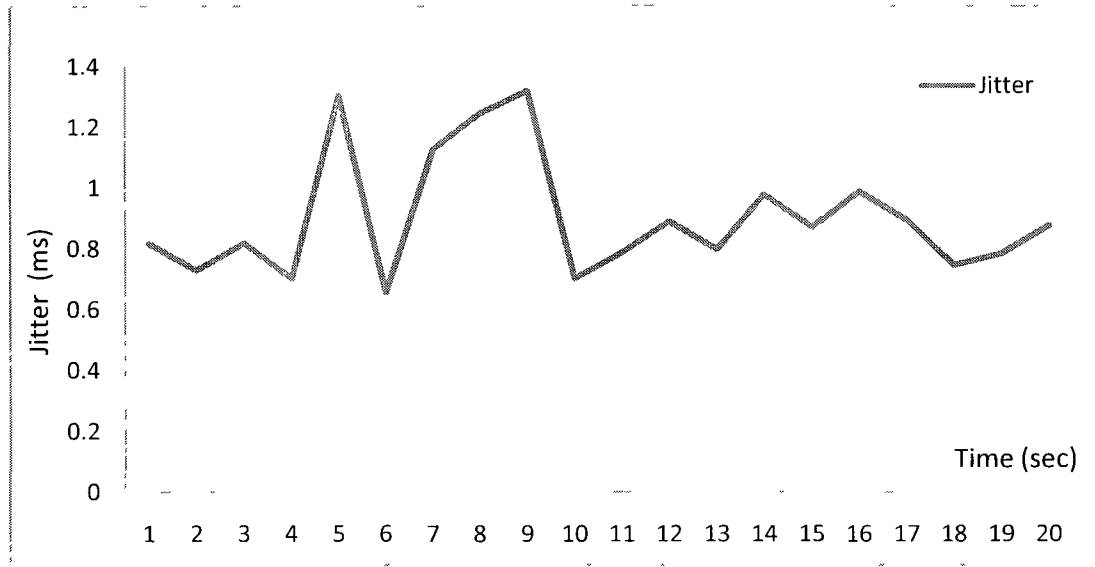


Figure 4.7: Jitter in two Paths Recovery Scenario

Figure 4.7 shows the delay jitter of packets in this scenario. It is noticed that the delay jitter hovers between 0.7 ms and 1.3 ms for the period from 1 to 20 seconds. In this case, the PLR detects the failure and then sends an RSVP message to the ingress router. Upon receiving the Error message, the ingress router redirects the traffic over the pre-provisioned backup LSP. This implies that the packets suffer a little delay jitter during the recovery process. Moreover, when the link failure occurs between the seconds 12 and 14, we observe a delay smaller than between seconds 7 and 10. This is due to the window size (8 Kbyte) that indicates the amount of data that can be buffered during a connection without a validation from the receiver.

4.2.3.3 Scenario#3 (MPLS Local Recovery)

In this scenario, a bidirectional dynamic path LSP is created between the two M10 and M160 LSR routers that traverse over on the POS OC-48 link. Furthermore, we deploy MPLS-TE FRR to automatically setup a detour LSP around the primary LSPs.

To simulate the link failure scenario, we shutdown the primary link interface, while, the secondary one is still connected. We repeat this step five times to get the accurate results.

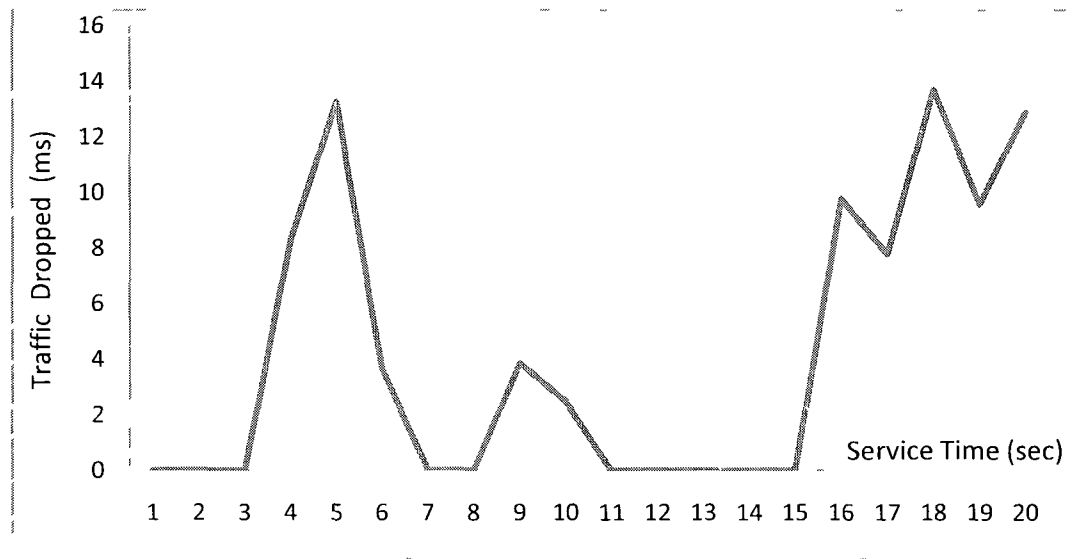


Figure 4.8: Traffic Dropped in MPLS Local Recovery Scenario

From Figure 4.8, we observe that the packet loss decreases from the second scenario to reach 14 ms in seconds 4 to 5, 18 to 19 and 10ms in second 16 resulting in 4.2 ms total packet loss. This indicates that the detour function by using MPLS-TE is very important to minimize the delay. Basically, PLR reroute the traffic on the detour link first and then send the RSVP Error messages to the ingress router to look for alternate new optimal path.

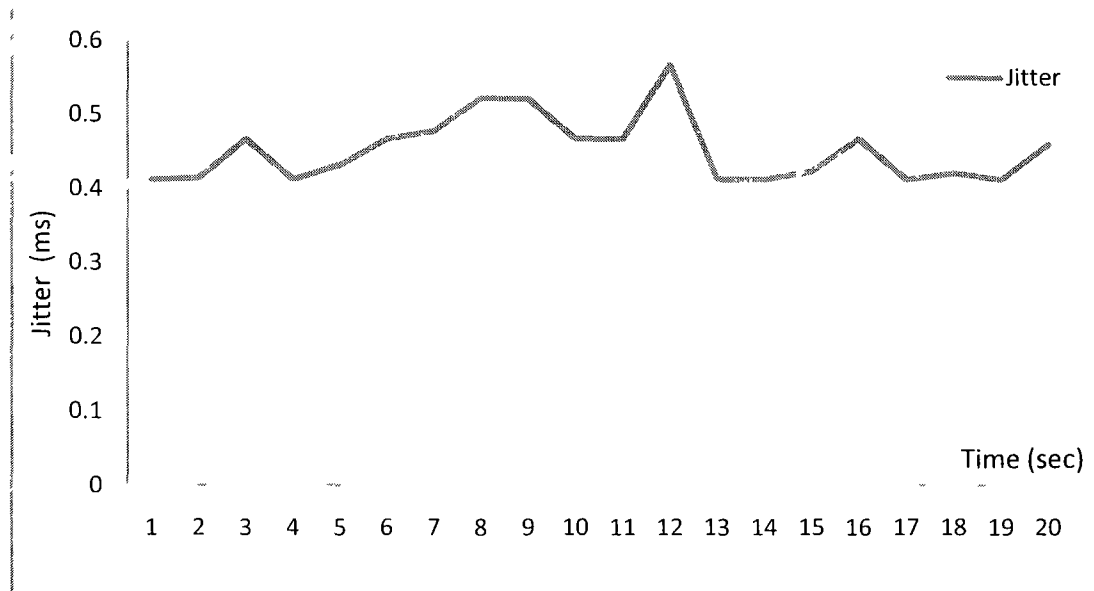


Figure 4.9: Jitter in MPLS Local Recovery Scenario

Figure 4.9 shows the delay jitter of packets in this scenario. It is observed that the delay jitter varies around 0.4 ms for the period from 1 to 20 seconds. In this case, once the PLR detects the failure, it immediately redirects the traffic over the pre-provisioned backup LSP. This implies that most packets traverse the same path after the switchover resulting in lower packet delay jitter. This can also be deduced comparing to results presented in Section 4.2.3.1 for the layer-3 routing convergence.

4.2.4 Network Recovery Result Comparison

Scenarios	Interval [sec]	File size [MB]	Bandwidth [Mbps]	Jitter [ms]	Packet loss/Total	Recovery Time [ms]
Layer-3 OSPF routing recovery	0 -120	1109	77.5	2.949	86777/877643	98.87
MPLS two paths recovery	0-60	478	66.8	0.921	3549/344291	10.30
MPLS local recovery	0-60	478	66.9	0.928	1448/342500	4.2

Table 4.1: Comparison of Different Network Recovery Schemes

Table 4.1 shows that the packet loss delay in the case of layer-3 OSPF routing recovery is the largest among the other schemes. However, the packet loss time decreases for the MPLS two paths recovery scheme to almost 10 ms, which indicates that it is improved by using MPLS based on RSVP signalling protocol. In addition, for the MPLS local recovery, the packet loss reduces to only 4.2 ms which proves that this scheme is the superior in terms of recovery time as

compared to the first two schemes because of using Fast Reroute mechanism. Furthermore, we notice that the MPLS global recovery scheme experiences a lower jitter value than that of the layer-3 OSPF routing recovery scheme. Moreover, the local recovery further decreases the jitter value to reach 0.4 ms due to the fast recovery by using the detour function while the other two schemes the packets suffer some delay.

Note that by knowing the rate of traffic sent (packet/sec) and the total loss of traffic due to the flap we are able to calculate the convergence time based on this equation:

Convergence time = Packet loss / Rate of traffic (packet/sec) (see Table 4.1).

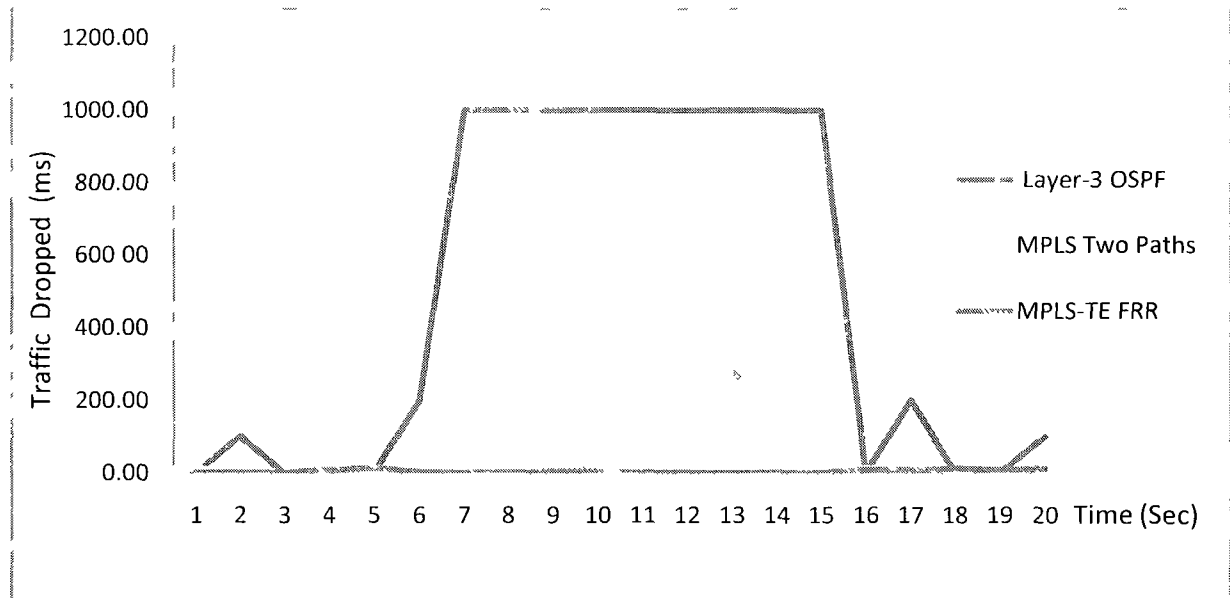


Figure 4.10: Traffic Dropped Comparison

Figure 4.10 shows that the packet loss is very large in the first scenario, which indicates that this scheme has a slow recovery process time which takes almost 9 seconds to get back to normal operation. On the other hand, in the second scenario, we notice that the packet loss delay is decreased to be 10.3 ms in seconds 11 to 13. However, the third scenario shows the best performance among the three scenarios, which is almost no packet dropped. This is due using the fast reroute mechanism.

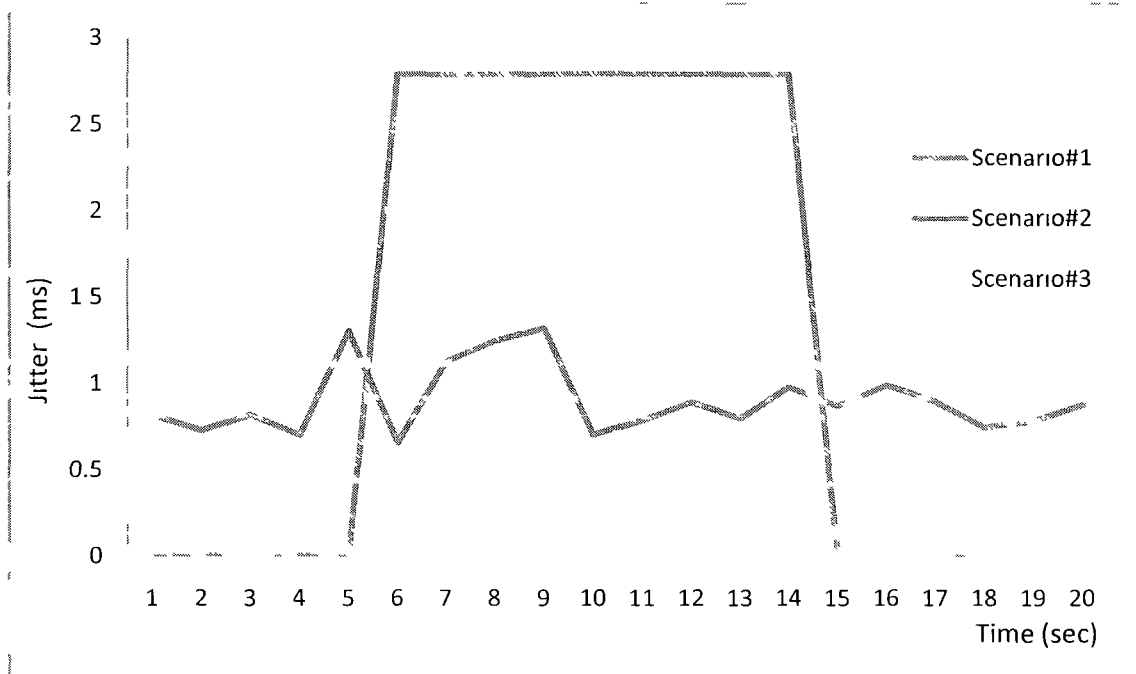


Figure 4.11: Jitter Comparison

From Figure 4.11 we observe that the jitter is improved in the second and third scenario. From both comparison charts, there is a clear indication that MPLS-TE FRR has significantly improved the performance and the convergence speed of the network.

4.3 Performance of Fast ReRoute Mechanisms

In this section, we compare the performance characteristics of packet loss and convergence time for two FRR mechanisms and their impact on IP and MPLS networks. Furthermore, we consider the effect of increasing a number of parameters on the overall convergence time and traffic loss, such as the number of LSPs affected by a single failure, and the number of routing prefixes.

In the first case, a MPLS-TE network is built with the consideration of FRR between the core routers along with the other routers. In the second case, a pure IP core network is built from the source to the destination, using different vendor routers with IP-FRR Loop Free Alternate (LFA) scheme as a recovery mechanism. Both IP-FRR and MPLS-TE FRR are used for performance evaluation in each testbed scenario implementation.

In order to carry out multiple traffic flows through our testbed, the Iperf traffic generator tool is also used to generate UDP traffic flows.

4.3.1 Scenario #4 (Number of LSPs versus Convergence time)

The provider's core transport infrastructure supports MPLS-TE. Our objective is to compute the convergence time and the packet loss, as well as study the effect of increasing the number LSPs between M10#2 and M160 that are affected by the link failure test

For that purpose, we create bidirectional dynamic Label Switched Paths (LSPs) between Juniper M10#2 and M160 router that cross on the primary GE link. For the sake of protection, we have deployed MPLS-TE FRR mechanism on M10#1 to automatically setup a detour around the primary LSPs as shown in Figure 4.12.

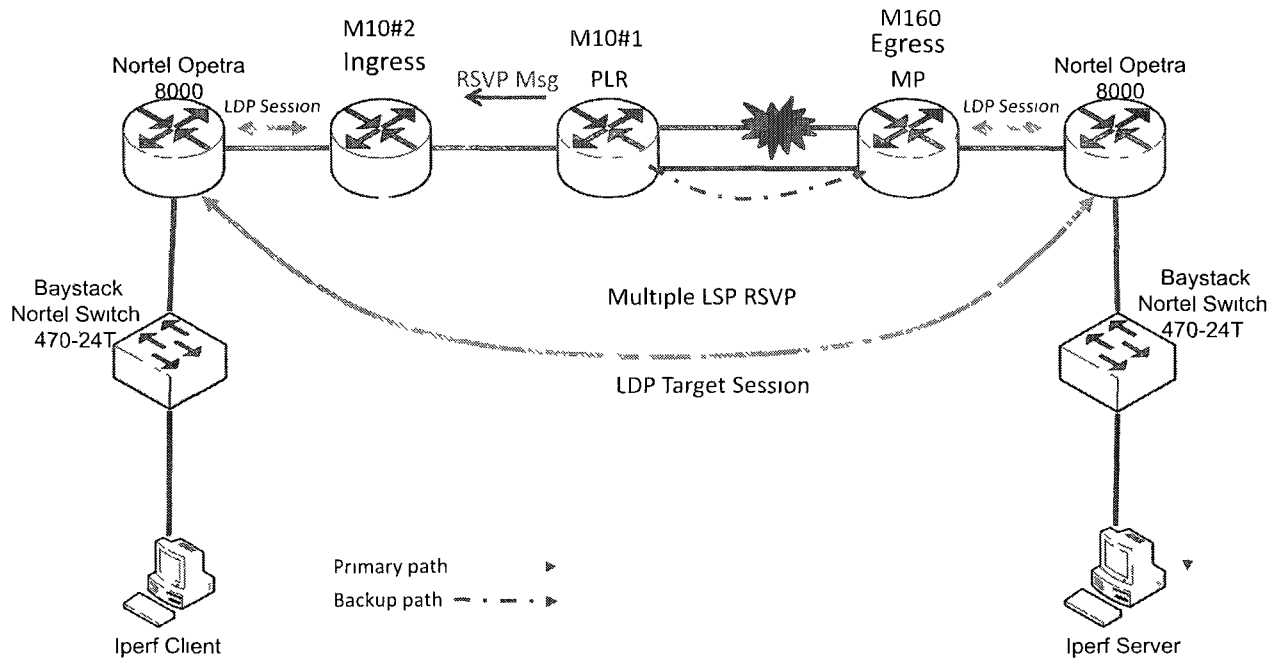


Figure 4.12: Link Failure in Multiple LSPs Scenario

As shown in Figure 4.12 above, the tests are performed by sending the traffic at the rate of 1000 packet/sec from the Iperf client to Iperf server at the far end through the MPLS core network. While the stream is being sent, the primary link between M10#1 and M160 is disconnected. Upon failure detection, the Point of Local Repair (PLR) reroutes the traffic over a pre-established backup LSP to maintain the flow of the protected traffic and guarantee minimal traffic loss. It then signals the outage to the ingress router M10#2 to notify the ingress router of the failure event by sending RSVP Path Error (PERR) messages that contains the address of the failed resource. Subsequently, the ingress router can exclude the failed link and compute and signal an optimal LSP over the remaining TE topology.

From the first scenario, with only one protected LSP, we can observe that the most relevant contribution on the overall convergence time (Calculated to be 70 ms) is the time needed to send notification messages from PLR (M10#1) to Ingress LSR (M10#2) and the time requested for Ingress LSR to process CSPF and signal a new optimal LSP as shown in Figure 4.13.

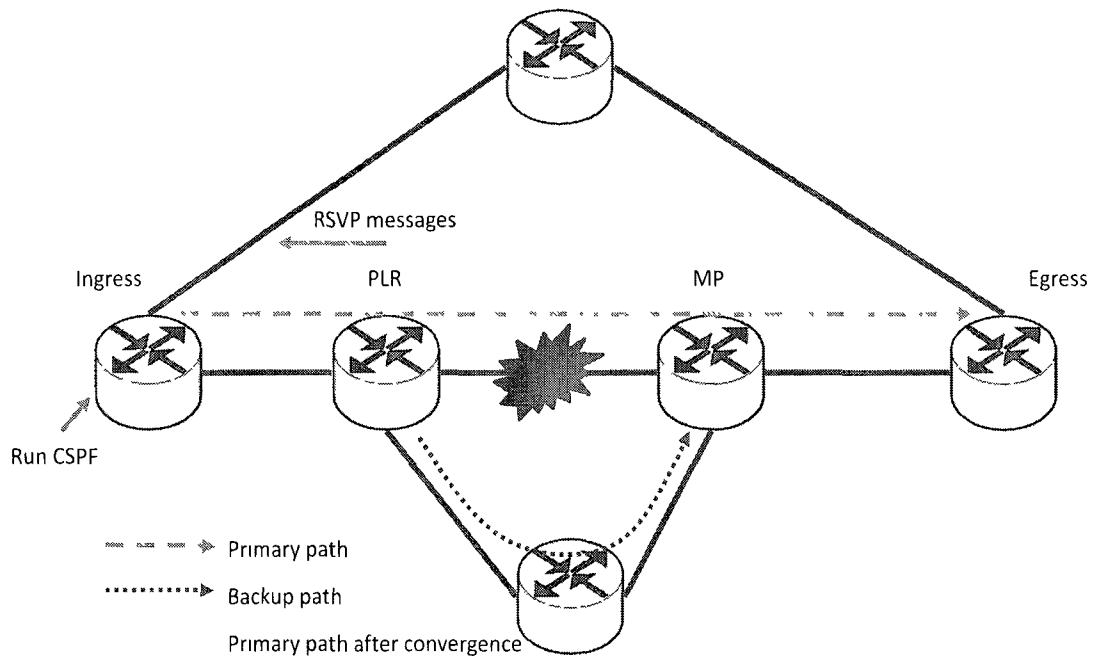


Figure 4.13: Link Failure in one Protected LSP Scenario

After the first scenario, in order to see the affect of increasing the number of LSPs on the convergence time, we repeat the same steps with different LSP numbers between M10#1 (Ingress) and M160 (Egress), using 10 to 300.

We describe the distribution of the convergence times for different LSPs in the network. Upon failure, the affected LSPs take varying convergence times to complete recovery after the fault onto the optimal path post the failure. This is accomplished by sending multiple RSVP Path Error (PERR) messages from PLR to the ingress router in the case of failure. In this case advisable differences could be noted for the recovery operation time; indeed a longer time is needed by the Ingress router to handle all notification messages, compute, and signal a new optimal LSPs.

The following figures show the distribution of convergence times for different LSP numbers.

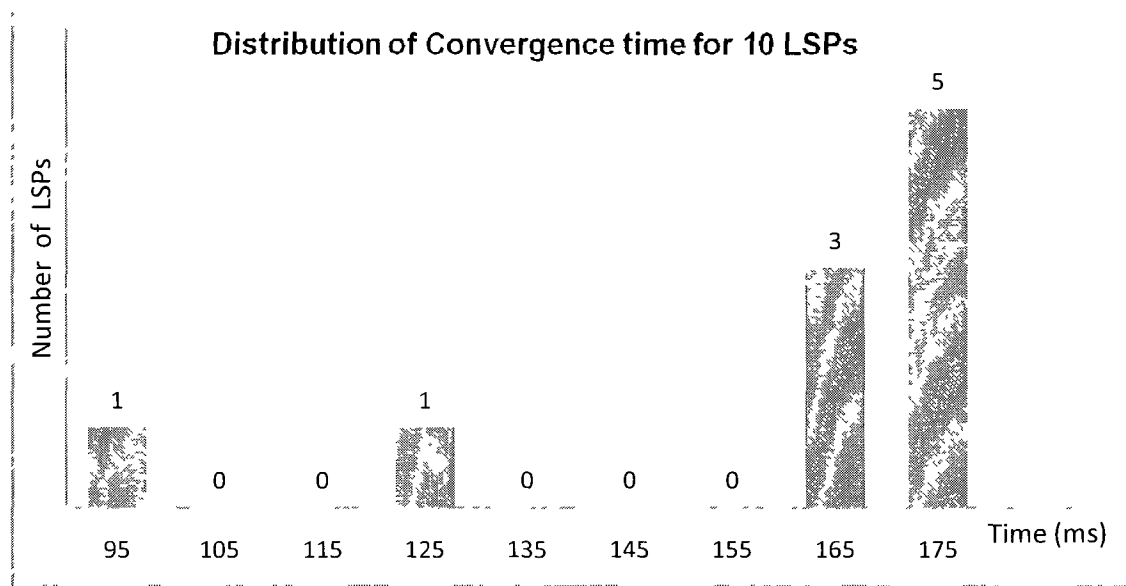


Figure 4.14: Distribution of 10 LSPs

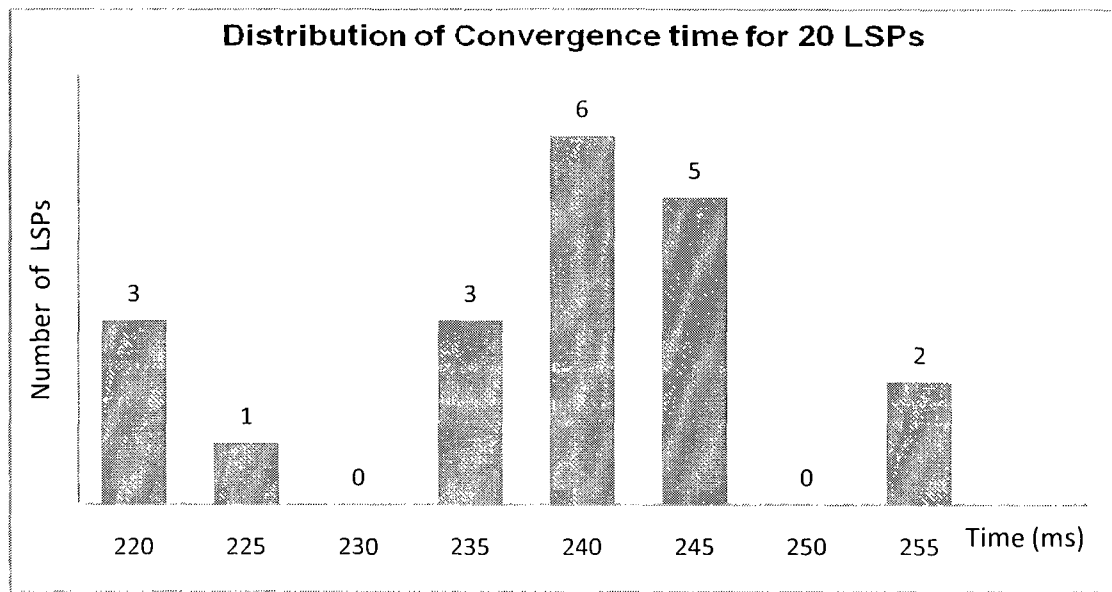


Figure 4.15: Distribution of 20 LSPs

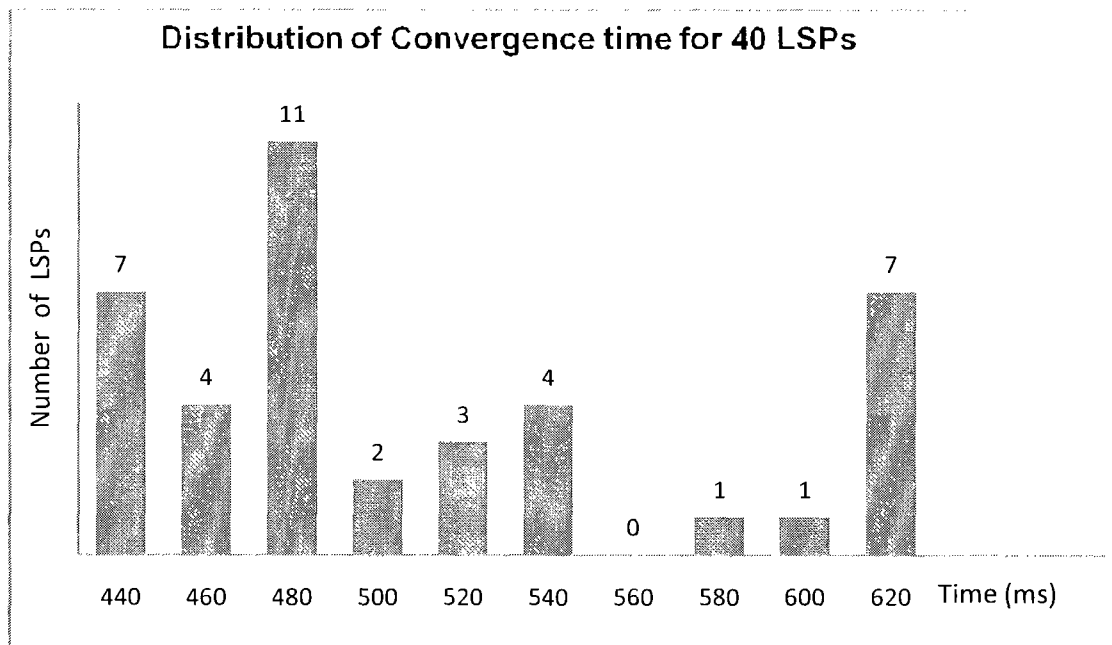


Figure 4.16: Distribution of 40 LSPs

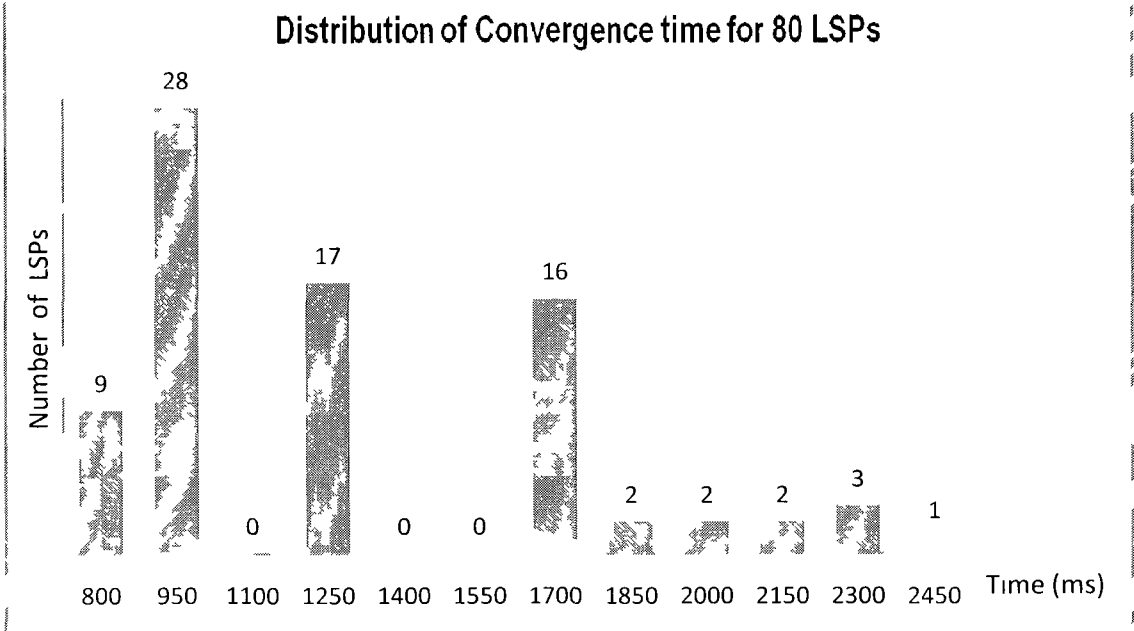


Figure 4.17: Distribution of 80 LSPs

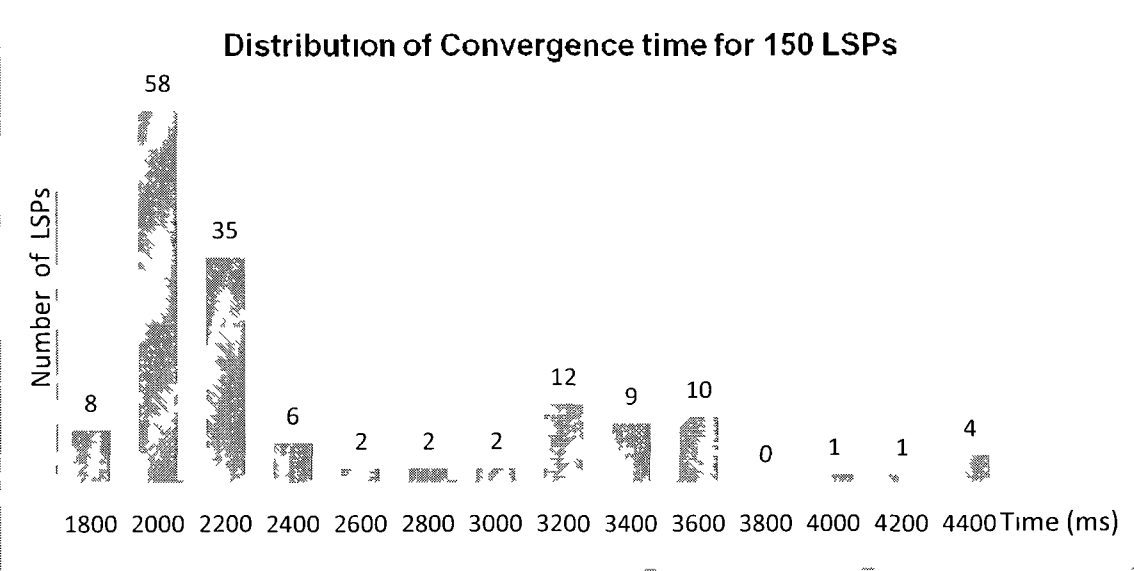


Figure 4.18: Distribution of 150 LSPs

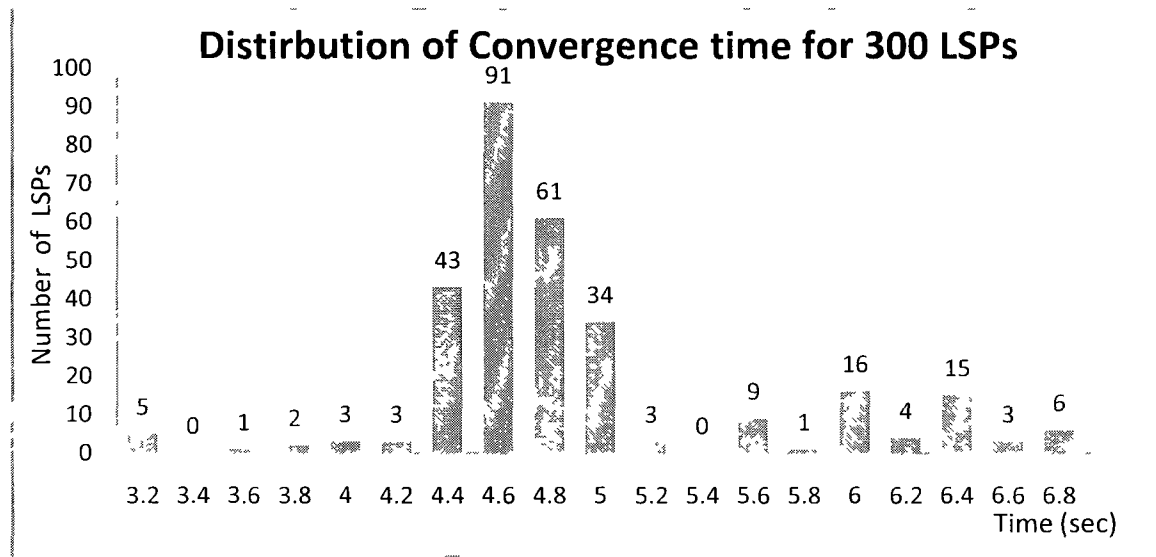


Figure 4.19: Distribution of 300 LSPs

From Figures 4.14 to 4.19, we observe that the recovery process can take several seconds depending on the number of tunnels or LSPs affected by the failure. For example, in Figure 4.16, the time needed for 11 LSPs is 480 ms and it goes up to 620 ms for the last 7 LSPs in the recovery process. In addition, from Figure 4.19, we observe that a longer convergence time is needed (3.2 sec to 6.8 sec) by the ingress router after the reception of the flooded notification messages for 300 LSPs which cause high routing load on processors to calculate and signal new optimal LSPs.

In each scenario, the numbers of packets lost due to the reroute to the other link are measured and the graph shown in Figure 4.20 is plotted.

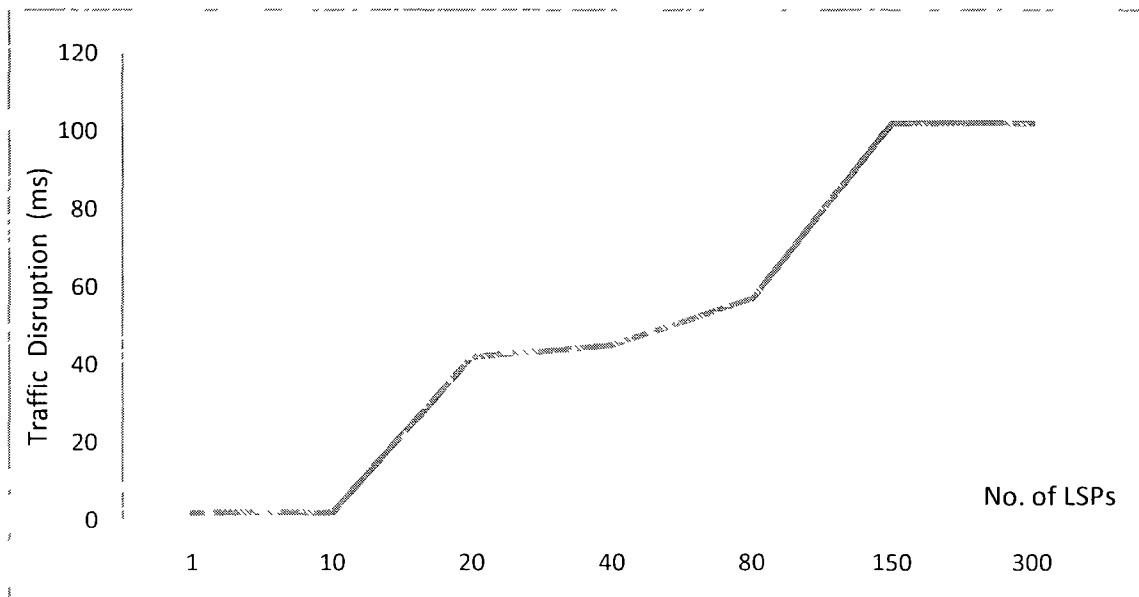


Figure 4.20: Traffic Disruption in Multiple LSPs Scenario

Figure 4.20 shows a higher packet loss is observed when increasing the number of LSPs affected by the link failure. For example, when 150 and 300 LSPs are configured to pass through the protected link, about 100 packets are lost from 1000 packets that are sent from the source to the destination. This means that about 100 ms is calculated as traffic loss time which is very high as compared to the 50 ms recovery time requirement by most real-time applications. This can be attributed to the PLR getting busier trying to switch more LSPs over to the backup paths post the link failure.

Likewise, the time it takes for the network to reroute to the preferred network link is measured. All the data is collected and the average convergence time in each scenario is calculated. A 95% confidence interval for the convergence time is plotted to give credibility, as shown below in Figure 4.21. The confidence intervals are not shown in the figure because the intervals are relatively small as seen in Appendix D.

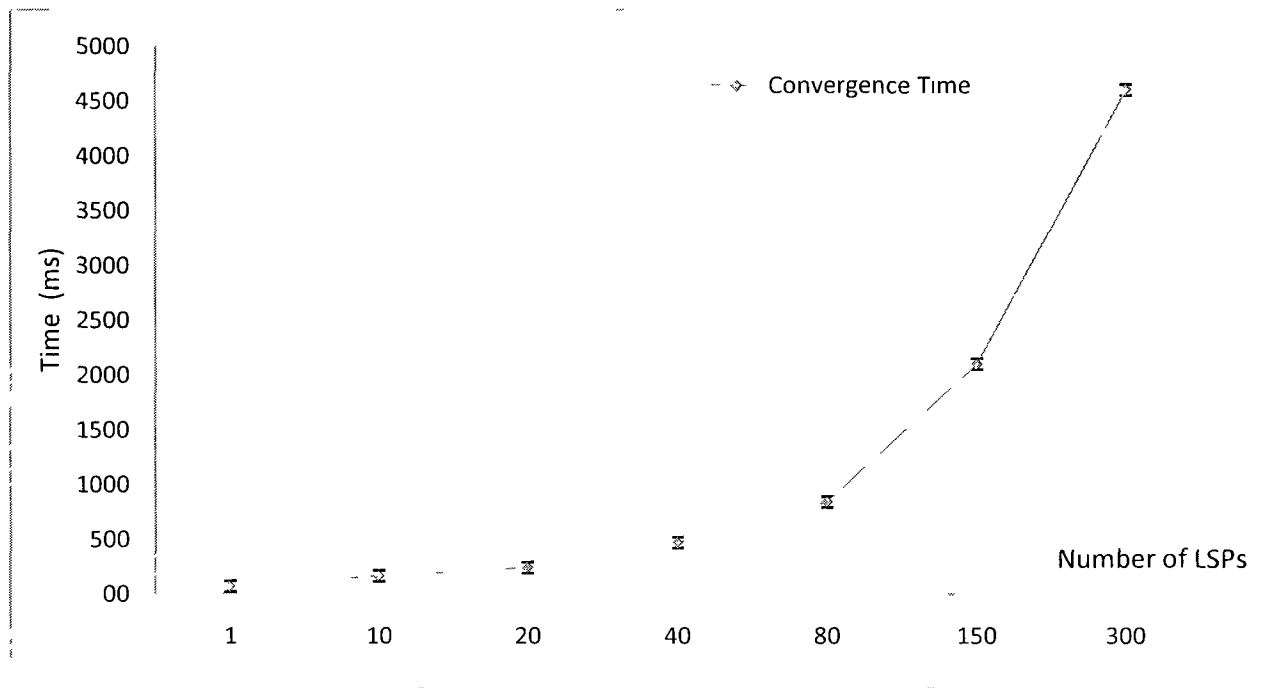


Figure 4.21: Number of LSPs vs. Convergence Time with Confidence Interval

As we can see from Figure 4.21, the convergence time increases as the number of LSPs increase, which makes it clear that RSVP has a scalability issue, when there are a large number of LSPs passing through the failed link. This is due to the amount of signalling that happens post the failure event (e.g. a PERR per affected LSP from PLR back to head-end) and the signalling of the new re-optimized LSP per affected LSP. Note that the PERR message is intercepted by each RSVP-TE LSR upstream of the PLR. In addition, the soft state characteristics of RSVP and RSVP-TE stresses the individual LSRs along the LSP paths due to the periodic refresh of the state for each LSP.

We conclude that in order to achieve acceptable convergence and switchover times, the MPLS network has to be enabled for TE-FRR and architected such that the number of LSPs or tunnels that are affected by any single failure is limited or minimized.

4.3.2 Scenario#5 (Number of Hops versus Convergence time)

In this scenario, our objective is to study the effect of increasing the number of hops between the ingress LSR and the PLR in an MPLS-TE FRR on the convergence time and the packet loss.

For this purpose, we have created bidirectional dynamic Label Switched Path (LSP) between Juniper M10#2 and M160 router that traverses the primary Gigabit Ethernet link. For the sake of protection, we have deployed MPLS-TE FRR mechanism on M10#2 to automatically setup a detour around the primary LSPs.

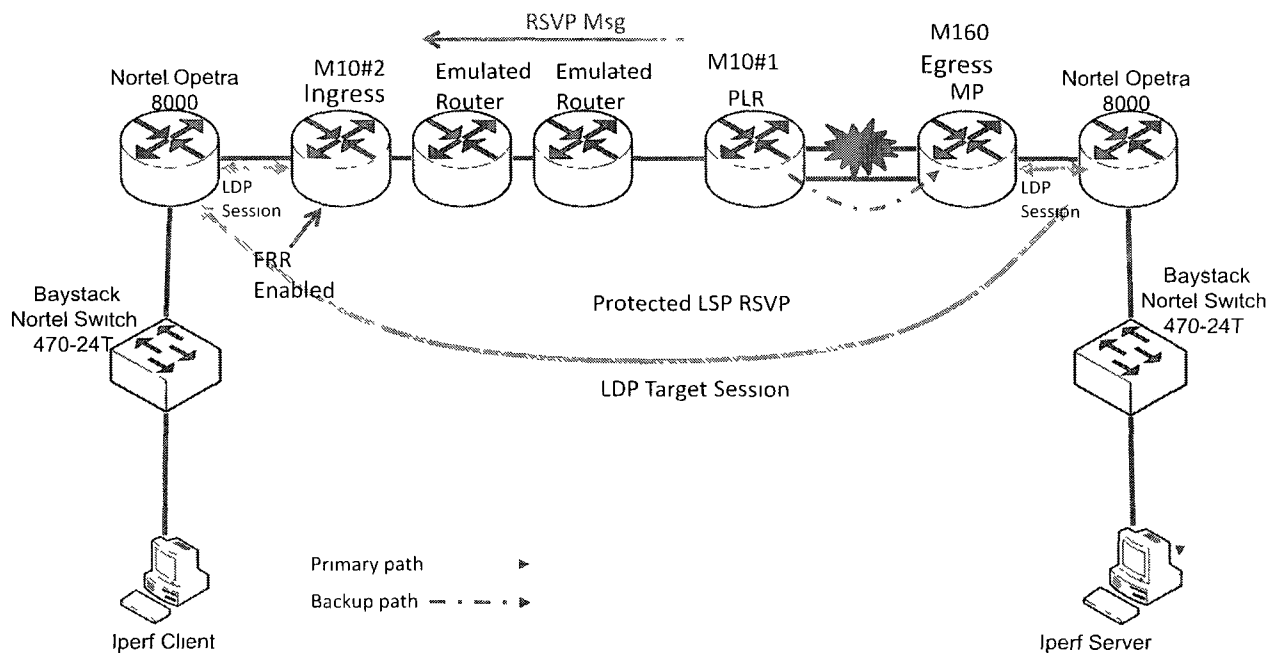


Figure 4.22: Link Failure in Multiple Hops Scenario

As shown in Figure 4.22 above, the tests are performed by sending the traffic at the rate of 100 packet/sec from the Iperf client to Iperf server applications running on PCs at each end through the core network. While the traffic is flowing, the used primary link between M10#1 and M160 is disconnected (by pulling the fiber). Upon failure detection, the PLR reroutes the traffic over a pre-established backup LSP to maintain the flow of the protected traffic and guarantees minimal loss. It then signals the outage to the ingress router M10#2 to notify the ingress router of the failure event. Subsequently, the ingress router can exclude the failed link and compute and signal an optimal LSP over the remaining TE topology.

From the first scenario, with 1 hop between ingress router and PLR, we observe that the most relevant contribution on the overall convergence time is the time needed to send notification messages from PLR (M10#2) to Ingress LSR (M10#1) and the time requested by the Ingress router to process CSPF and signal a new optimal LSP.

After the first scenario, we have repeated the same steps with different hops numbers between M10#2 (Ingress) and M10#1 (PLR), using 2 to 5.

In each scenario, the numbers of packets lost due to the reroute to the backup link are measured and the graph shown in Figure 4.23 is plotted.

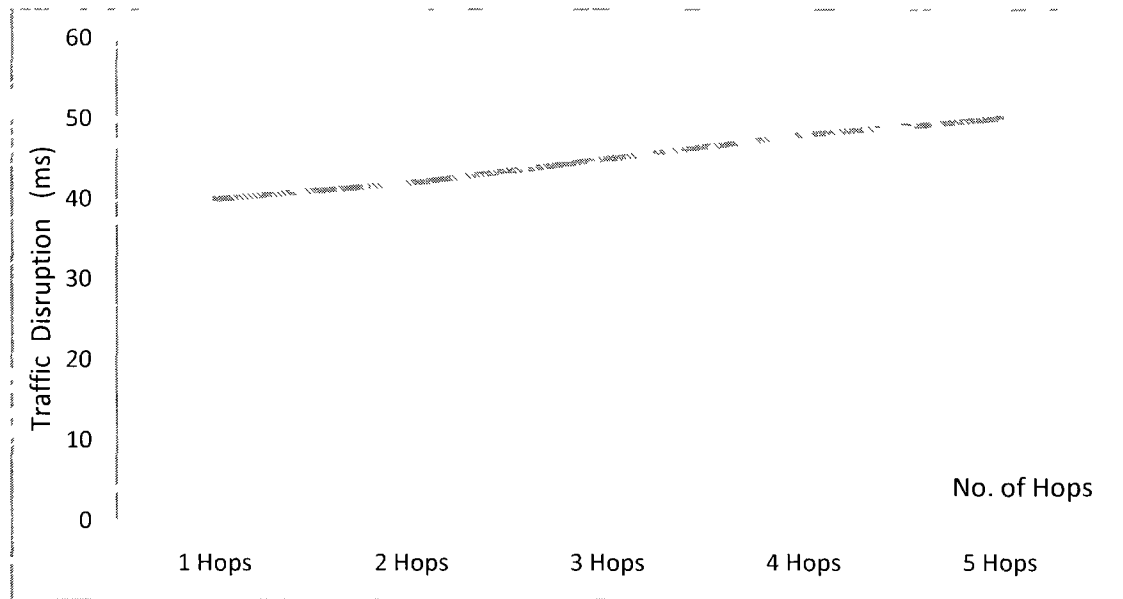


Figure 4.23: Traffic Disruption in Multiple Hops Scenario

From Figure 4.23, we observe that the traffic loss is almost un-affected when increasing the number of hops between the ingress router and PLR. This indicates that the signalling delay per LSP hop is negligible as compared to the overall convergence time during the recovery process. For example, only 4 packets are lost when sending 100 packets in all the scenarios. This means about 40 ms is recorded as a traffic loss time.

Likewise, the time it takes for the network to reroute to the preferred network link is measured. The average convergence time in each scenario is calculated. A 95% confidence interval for the convergence time is plotted to give credibility, as shown below in Figure 4.24. The confidence

intervals are not shown in the figure because the intervals are relatively small as seen in Appendix D.

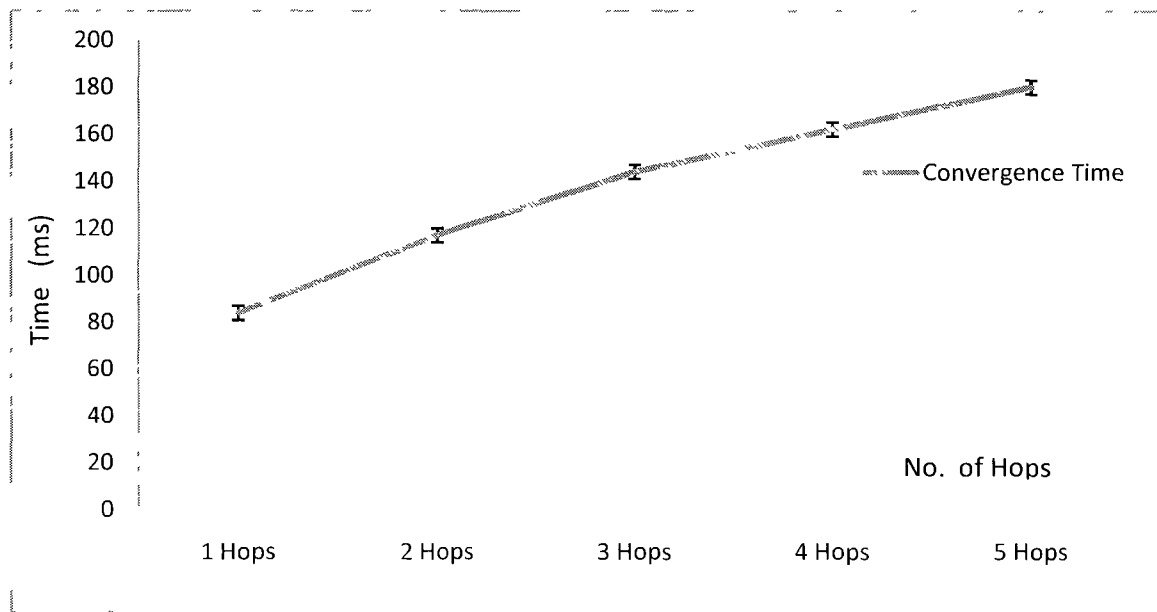


Figure 4.24: Number of Hops vs. Convergence Time with Confidence Interval

We observe from Figure 4.24, the convergence time slightly increases as the number of hops increases. This increase is due to the RSVP Path error message taking longer times to travel along the hops from the PLR router to the ingress router. In other words, each hop adds some delay to the RSVP message which increases the overall convergence time.

4.3.3 Scenario#6 (Number of IGP Prefixes versus Convergence time)

The provider's core transport infrastructure supports pure IP with Fast ReRoute protection. Our objective in these tests is to compute the convergence time and the packet loss and study the effect of increasing the IGP prefixes injected into the network.

For that purpose, we have created a ring network as shown in Figure 4.25 below and have enabled the Fast ReRoute mechanism on M10#2 which chooses the Emulated router as a primary next-hop (E), based on its lower cost, and M10#1 as a second next-hop (N-1) to protect the network in case of failure. Both of the nodes appear in the routing table of M10#2, one as a primary destination and the second as a backup destination (see details in Appendix C).

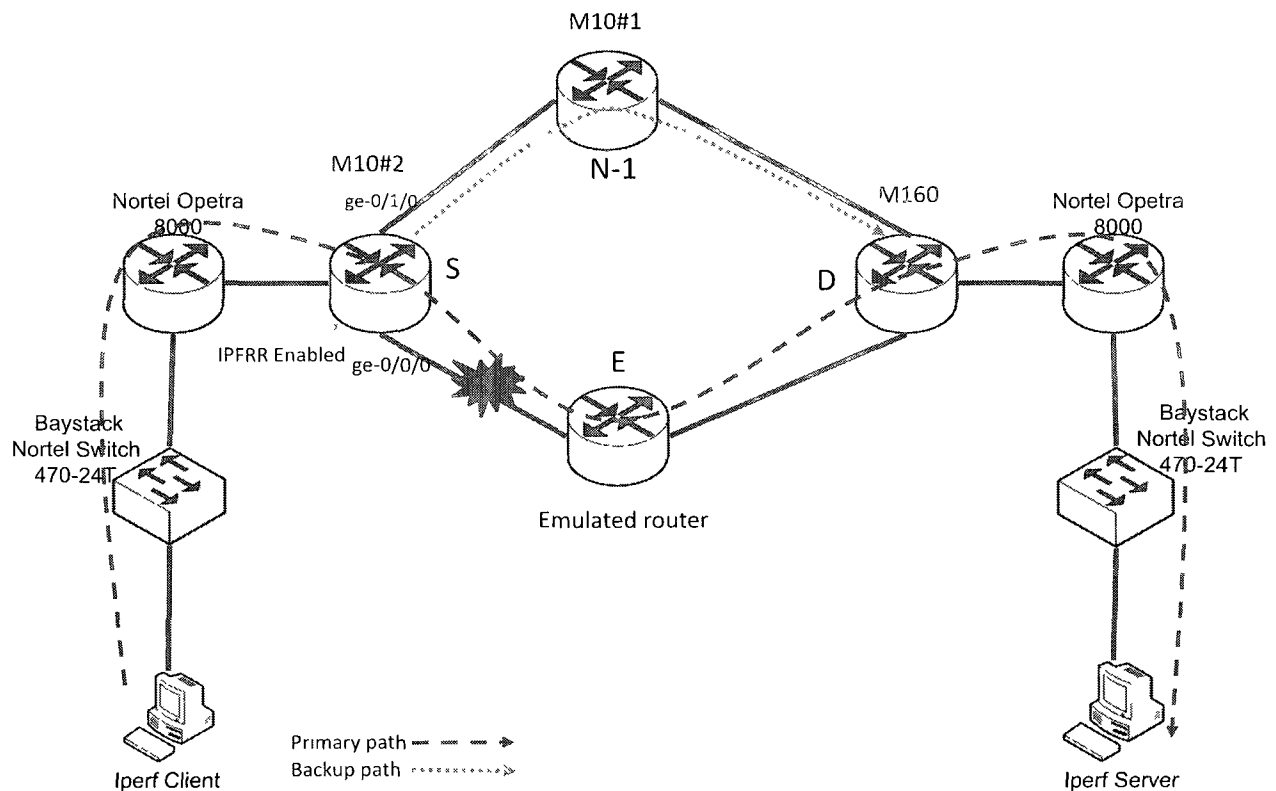


Figure 4.25: Link Failure in IP-FRR Scenario

In this scenario, the tests have been performed by sending the traffic at the rate of 500 Packets/sec from the Iperf client to the Iperf server at the far end through the core network. When the stream is being sent, the used primary link between M10#2 and emulated router (E) is disconnected.

Upon failure detection, M10#2 reroutes the traffic over the backup next-hop in its routing table, which is M10#1, to maintain the flow of the protected traffic and guarantee a faster recovery. After the first scenario, we repeat the same steps by increasing the number of IGP prefixes announced to be appeared in the routing table, using numbers 10 to 200.

We repeated the tests three times for each scenario and collected all the data. From the data, we calculated the average of the convergence time and the standard deviation for each scenario. We then calculated 95% confidence interval for each scenario as shown in Table D.1. The average of the convergence times for the tests is shown in Figure 4.26. The confidence intervals are not shown in the figure because the intervals are relatively small as seen in Appendix D.

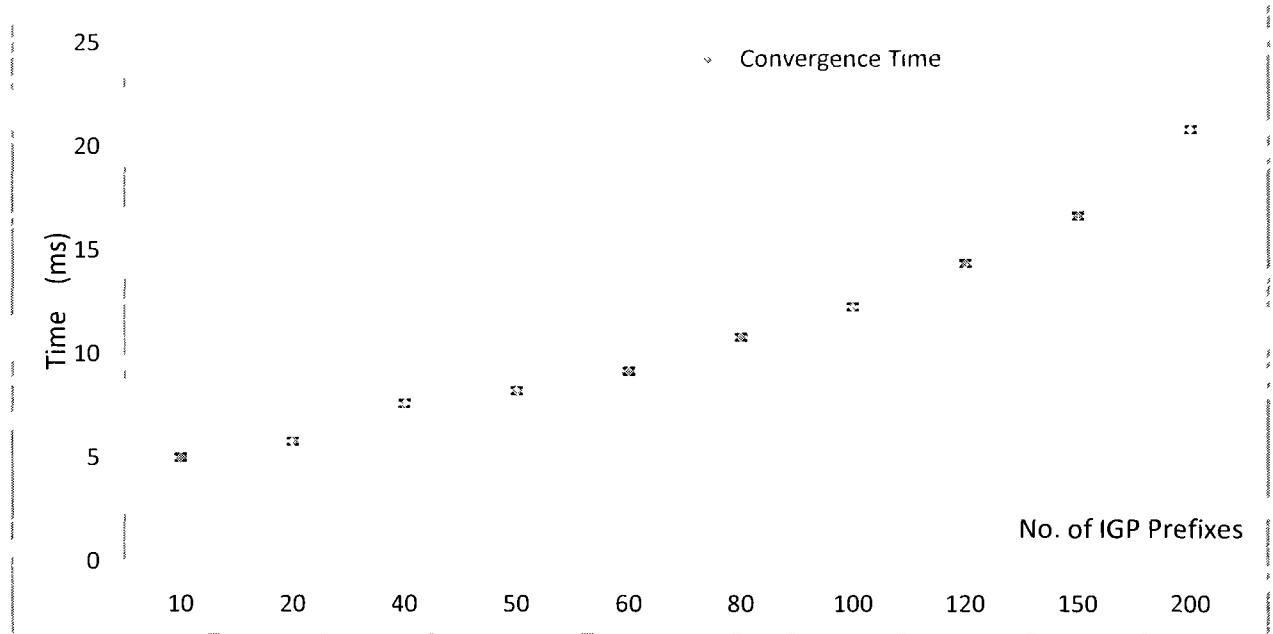


Figure 4.26: No. of IGP Prefixes vs. Convergence Time with Confidence Interval

Figure 4.26 shows that the convergence time is not linear when the number of IGP prefixes is increased. Specifically, it goes up from 5 ms in 10 prefixes to reach 20 ms in 200 prefixes case. As a result, we attribute this to the PLR getting busy trying to update the next-hop per-prefix in its routing table. In other words, PLR takes a longer time to setup new entries after the link failure by starting a new SPF process and new alternate entries by starting a new backup SPF process (see details in Appendix C).

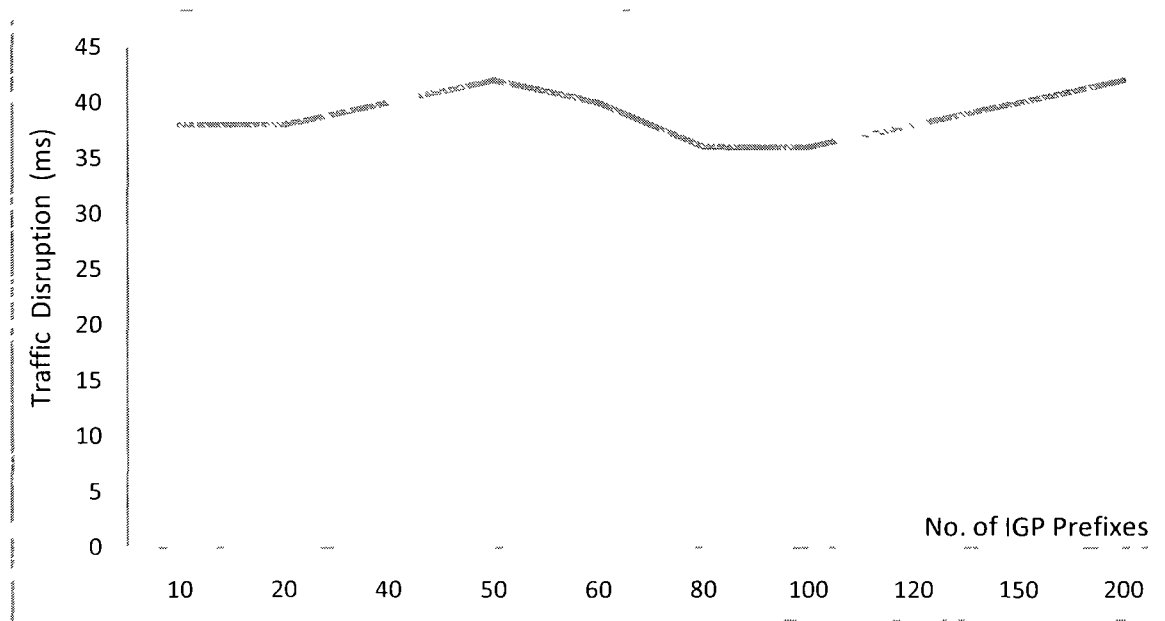


Figure 4.27: Traffic Disruption in IP-FRR Scenario

From Figure 4.27, we observe that the number of packets lost is almost unaffected when increasing the number of IGP prefixes. For example, in each scenario, we lost about 20 packets from sending about 500 packets in total, which is about 40 ms traffic loss time. This time is considered to be small time, which indicates that IP-FRR is a good solution for protecting the traffic from being lost and give competitive convergence time to SONET's convergence time (50 ms). Moreover, IP-FRR depends on IGP which is simpler than using RSVP for signalling.

4.4 Summary

Building MPLS systems that can survive network failures is not simple. This is due to the fact that MPLS is built on top of IP, which has less demanding recovery requirements and has its own method of resolving routing changes. It also owes something to the origins of MPLS, which, when originally designed, paid little attention to rapid recovery from failures.

This chapter presented our methods and tools, which allow us to evaluate the performance of network recovery mechanisms across multiple heterogeneous networks. These methods and tools are carried out by some experiments using testbed environments in our Lab. We have described the architecture for the testbed in each scenario and the methodology that has been used in each scenario.

The test results show that the local protection scheme (MPLS-TE FRR) is the best suitable technique with respect to packet loss and jitter in case of link failure, which comply with the results in [ALB05][ALA07]. However, the convergence time in the case of IP-FRR increases as the routing table increases by IGP prefixes, which also complies with what is published in [FRA05]. On the other hand, the convergence time in MPLS-TE FRR also increases as the number of protected LSPs in the core network increases which is also explained in Chapter 3 of [MIN08]. As well as, the number of hops between the ingress and PLR router increase the convergence time. Therefore, for any operator that wants to deploy any of these networks, he should consider these parameters.

Chapter 5

Concluding Remarks and Future Work

5.1 Concluding Remarks

Recovery from network failures is a very important and challenging concern for many Internet service providers. Real-time applications such as voice and video require networks to respond quickly to failures, and to return to normal operation in order to preserve the required quality.

There are two ways for networks to recover from failures and provide fast recovery time. The first is protection at the physical layer, which includes SONET and WDM. These technologies provide a less than 50 ms recovery time. However, they require dedicated hardware and resources which are very expensive.

The second is protection at the network layer, which deploys pure IP networks and relies on Interior Gateway Protocol (IGP) convergence for the protection. Even though relying on IGP convergence to overcome the failure has some advantages, such as, quickly routing around failures and converging on the remaining topology, it suffers the disadvantage of taking a few seconds (5-10 sec) to recover from a failure [ZAB06].

As a result, IP-FRR and MPLS-TE FRR are designed to overcome several well-known limitations of traditional IP networks, as well as fulfill the sensitive application needs by providing lower cost solutions than the physical approaches [GJO07] [CHU00].

In this thesis, we have provided methods and tools that study and evaluate the performance and capability of the network protection schemes in both IP and MPLS core networks. We have simulated a network in a testbed environment that consists of both real and emulated routers. Our results show that the packet loss and jitter is improved by using local protection at the network level.

In addition, we have developed methods and tools that present and evaluate the performance of Fast ReRoute mechanisms. These mechanisms give a competitive recovery time to SONET's recovery time. We also have investigated the performance of these techniques by implementing some different scenarios in our testbed. Furthermore, we have discussed how these schemes act when a network failure occurs and have given the advantages and disadvantages of each technique.

To conclude, in order to give an MPLS network the ability to converge faster, we should use MPLS-TE Fast ReRoute as a protection mechanism with the consideration of the number of hops between the ingress router and the PLR router as well as the number of LSPs passing through the protected link. For example, to minimize or limit the convergence time, the network operator can limit the number of hops between the ingress and PLR and also the number of protected LSPs.

Another solution is to deploy our network with IP-FRR LFA which is much simpler to deploy on pure IP and MPLS than MPLS-TE. LFA is a solution directly integrated into the IGP which is simpler when a universal coverage is needed than a scheme based on universal coverage delivered via RSVP. In addition, any one before designing a network that is based on IP-FRR should consider the number of IGP prefixes. For example, to minimize or limit the convergence time, the number of IGP prefixes should be limited.

5.2 Future Work

Although the tests in this thesis are comprehensive, much work in this area remains to be carried out. In our first three scenarios, we have used the default parameters of OSPF hello packets. Future studies might attempt to change the hello timers in order to see their affect on the convergence. Another possibility would be to test IS-IS or BGP instead of OSPF. This might give faster convergence in the access links, since, for example, IS-IS convergence can be accomplished in milliseconds in small topologies.

In this thesis, we have also evaluated Fast ReRoute schemes, i.e. MPLS-TE Fast ReRoute and Loop Free Alternate FRR (LFA FRR), by carrying out some tests in a testbed environment. We

have demonstrated the relation between the convergence time in MPLS-TE FRR and LSPs number. In future research, we suggest implementing a different IP Fast ReRoute technique and compare the results with LFA FRR.

This thesis has mainly looked at recovery from link failures and control plane failures, which can be detected by the RSVP-TE hello mechanism in MPLS-TE FRR or IGP hello packets in case of IP-FRR. Currently the Bidirectional Forwarding Detection (BFD) protocol is developed to implement a faster data plane failure detection mechanism. The field of failure detection in the data plane remains open to research.

The scenarios we have presented in this thesis look at link failures. Considering node failures would be another great contribution to the evaluation of the schemes discussed in this thesis.

References

- [ABD04] H. Abdalla , A.M. Soares, P.H.P. De Carvalho, G. Amvame-Nze, P. Solís Barreto, R. Lambert, E. Pastor, I. Amaral, V. Macedo , and P. Tarchetti, “Performance Evaluation of Shortest Path Computation for IP and MPLS Multi-service Networks over Open Source Implementation”, Proceeding of SAPIR Conference, Springer-Verlag, Berlin, pp. 267-278, 2004.
- [ALA07] B. Alawieh, R. Ahmed, and H. Mouftah, “Measurement of Restoration Time from Single-Link Failure in an Optical testbed Environment”, Proceedings of 3rd International Conference on Testbeds and Research Infrastructure for the Development of Networks and Communities, TridentCom, Lake Buena Vista, FL, USA, May 2007.
- [ALB05] R. Albanese, D. Ali, S. Giordano, U. Monaco, F. Mustacchio , and G. Procissi, “Experimental Comparison of Fault Notification and LSP Recovery Mechanisms in MPLS Operational Testbeds”, Proceeding of QoS-IP, Springer-Verlag, Berlin Heidelberg, pp. 136–149, 2005.
- [ANE99] P. Anelli, and M. Soto, “Evaluation of the APS protocol for SDH rings reconfiguration”, IEEE Transactions on Communications, vol. 47, pp. 1386-1393, Sep 1999.
- [ASI06] H. M. Asif, and Md. G. Kaosar, “Performance Comparison of IP, MPLS and ATM Based Network Cores using OPNET”, IEEE, Source: Institute of Electrical and Electronics Engineers , Aug 2006.
- [ATL08] A. Atlas, and A. Zinin, “Basic Specification for IP Fast Reroute: Loop-Free Alternates”, IETF, RFC 5286, Sep 2008.

- [AWD99a] D. O. Awduche, "MPLS and Traffic Engineering in IP Networks", IEEE communication magazine, vol. 37, pp 42-47, Dec 1999.
- [AWD99b] D. Awduche, J. Malcolm, J. Agogbua, M. O'Dell and J. McManus, "Requirements for Traffic Engineering Over MPLS", IETF, RFC 2702, Sep 1999.
- [BAS01] A. Basu and J. G. Riecke, "Stability Issues in OSPF Routing", Proceedings of SIGCOMM 01, Special Interest Group on data Communications, Aug 2001.
- [BRA97] R. Braden, L. Zhang, S. Berson, S. Herzog, and S. Jamin, "Resource Reservation Protocol (RSVP)", ver.1 Functional Specification, RFC 2205, Sep 1997.
- [BRY07] S. Bryant, M. Shand, C. Filsfils, and S. Previdi, "IP Fast ReRoute using tunnels, draft-bryant-ipfir-tunnels-03.txt", Nov 2007 work in progress.
- [CHA02] H. Changcheng, V. Sharma, K. Owens, and S. Makam, "Building reliable MPLS networks using a path protection mechanism", IEEE Communications Magazine, pp. 156-162 2002.
- [CHU00] J. M. Chung, "Analysis of mpls traffic engineering in Circuits and Systems", Proceedings of the 43rd IEEE, Midwest Symposium, vol. 2, pp. 550-553, 2000.
- [CLA88] D. D. Clark, "The Design Philosophy of the DARPA Internet Protocols", Proceedings of SIGCOMM 88 on Computer Communication, vol. 18, pp. 106-114, Aug 1988.
- [DAS02] S.K. Das, P. Venkataram, "A method of designing a path restoration scheme for MPLS based network", Proceeding of 5th IEEE International Conference on High Speed Networks and Multimedia Communications, 2002.

[**ENY09**] G. Enyedi, G. Rétvári, P. Szilágyi, and A. Császár, “IP Fast ReRoute: Lightweight Not-Via”, Proceeding of IFIP International Federation for Information, pp. 157–168, 2009.

[**FOO03**] J. Foo, “A Survey of Service Restoration Techniques in MPLS Networks”, 2003 available online at <http://citeseerx.ist.psu.edu/viewdoc/summary?doi=10.1.1.6.3402>, visited in Sep 2010.

[**FRA05a**] P. Francois, C. Filsfils, J. Evans, and O. Bonaventure, “Achieving Sub-Second IGP Convergence in Large IP Networks”, Proceedings of ACM SIGCOMM Computer Communication, vol. 35, 2005.

[**FRA05b**] P. Francois and O. Bonaventure, “An evaluation of IP-based Fast ReRoute techniques”, Proceedings of CoNext, Toulouse, France, 2005.

[**GER99**] O. Gerstel, P. Lin, and G. Sasaki, “Combined WDM and SONET Network Design”, Proceedings of INFOCOM Eighteenth Annual Joint Conference of the IEEE Computer and Communications Societies, vol.2, pp. 734 – 743, Mar 1999.

[**GJO07**] M. Gjoka, V. Ram, and X. Yang, “Evaluation of IP Fast ReRoute Proposals”, Proceedings of 2nd International Conference on Communication Systems Software and Middleware, COMSWARE, 2007.

[**GOY03**] M. Goyal, K. K. Ramakrishna, and W. Fang, “Achieving Faster Failure Detection in OSPF Networks”, Proceedings of IEEE International Conference on Communications, vol.1, pp 296 – 300, 2003.

[**GRE09**] H. Gredler, “50ms Protection for IP/LDP Networks”, Presentation from Juniper Networks: Presented on Nov 2009. Available online at: <http://www.juniper.net/us/en/community/junos/live/091103/>, Visited in Sep 2010.

[HUA07a] “MPLS TE Fast ReRoute”, Technical White Paper, Huawei Technologies Co., Ltd. 2007.

[HUA07b] “LDP over RSVP”, Technical White Paper, Huawei Technologies, 2007.

[HYU00] H. Yu, C. Alaettinoglu, and V. Jacobson, “Towards Milli-Second IGP Convergence”, IETF Internet Draft: draft-alaettinoglu-ISIS-convergence-00, Nov 2000.

[JAM02] B. Jamoussi, L. Andersson, R. Callon, R. Dantu, L. Wu, P. Doolan, T. Worster, N. Feldman, A. Fredette, M. Girish, E. Gray, J. Heinanen, T. Kilty, and A. Malis, “Constraint-Based LSP Setup using LDP”, IETF, RFC 3212, Jan 2002.

[JIA09] X. Jiang, M. Xu, Q. Li, and L. Pan “Improving IGP Convergence through Distributed OSPF in Scalable Router”, Proceedings of 11th IEEE International Conference on High Performance Computing and Communications, pp 438-443, Jun 2009.

[JUN10] “Ensuring Rapid Restoration in Junos OS-Based Networks”, Technical white paper, Juniper Networks, 2010.

[KAT10a] D. Katz, and D. Ward, “Bidirectional Forwarding Detection (BFD)”, Internet Engineering Task Force (IETF), RFC 5880, ISSN: 2070-1721, Jun 2010.

[KAT10b] D. Katz, and D. Ward, “Bidirectional Forwarding Detection (BFD) for IPv4 and IPv6 (Single Hop)”, Internet Engineering Task Force, RFC 5881, Jun 2010.

[LAK05] U. Lakshman, and L. Lobo, “MPLS Configuration on Cisco IOS Software”, Series: Networking Technology, ISBN-13: 978-1-58705-199-9, Oct 2005.

[MEN10] M. Mentha, M. Hartmann, R. Martina, T. Čičić, and A. Kvalbeinc, “Loop-free alternates and not-via addresses: A proper combination for IP Fast ReRoute?”, Resilient and Survivable networks, vol. 54, pp 1300-1315, Jun 2010.

- [MET02] C. Metz, “IP protection and restoration”, Proceeding of Internet Computing, IEEE vol. 4, Issue: 2, pp 97 – 102, Apr 2000.
- [MIN08] I. Minei and J. Lucek, “MPLS-enabled application: emerging developments and new technologies 2nd Edition”, ISBN: 978-0-470-98644-8, May 2008.
- [MOY98] J. Moy, “OSPF Version 2”, IETF, RFC 2328, Apr 1998.
- [NAD10] T.D. Nadeau, A.K. Koushik, and R. Cetin, “Multiprotocol Label Switching (MPLS) Traffic Engineering Management Information Base for Fast ReRoute”, Internet Draft, Apr 2010.
- [PAN05] P. Pan, G. Swallow, and A. Atlas, “Fast ReRoute Extensions to RSVP-TE for LSP Tunnels”, IETF, RFC 4090, May 2005.
- [PAS08] S. Pasqualini, A. Iselt, A. Kirstadter, and A. Frot, “MPLS Protection Switching vs. OSPF Rerouting A Simulative Comparison”, Source:Siemens, Jan 2008.
- [RAH07] M. Rahmani, W. Hintermaier, B. M`uller-Rathgeber, and E. Steinbach, “Error Detection Capabilities of Automotive Network Technologies and Ethernet - A Comparative Study”, Proceedings of the IEEE Intelligent Vehicles Symposium, Istanbul, Turkey, Jun 2007.
- [RAH08] M. A. Rahman, A. H. Kabir, K. A. M. Lutfullah, M. Z. Hassan, and M. R. Amin, “Performance Analysis and the Study of the behaviour of MPLS Protocols”, Proceedings of the International Conference on Computer and Communication Engineering, Kuala Lumpur, Malaysia, pp. 226 – 229, May 2008.
- [RAJ07] A. Raj, and O. C. Ibe, “A survey of IP and multiprotocol label switching Fast ReRoute schemes”, Computer Networks, vol. 51, pp. 1882-1907, Jun 2007.
- [ROS01] E. Rosen, A. Viswanathan, and R. Callon, “Multiprotocol Label Switching Architecture”, IETF, RFC 3031, Jan 2001.

- [SAN01] Y. Sangsik, L. Hyunseok, C. Deokjai, K. Youngcheol, L. Gueesang, and M. Lee, “An efficient recovery mechanism for MPLS-based protection LSP”, Proceedings of Joint 4th IEEE International Conference on ATM and High Speed Intelligent Internet Symposium, 2001.
- [SCH03] G. Schollmeier, J. Charzinski, A. Kirstadter, C. Reichert, K. J. Schrodi, Y. Glickman, and C. Winkler, “Improving the Resilience in IP Networks”, Proceedings of IEEE Conferences on High Performance Switching and Routing, pp 91 – 96, 2003.
- [SHA03] V. Sharma, and F. Hellstrand, “Framework for multi-protocol label switching (MPLS)-based recovery”, IETF RFC 3469, Feb 2003.
- [SHA10a] M. Shand, and S. Bryant, “IP Fast ReRoute Framework”, Internet Engineering Task Force (IETF), RFC 5714, Jan 2010.
- [SHA10b] M. Shand, S. Bryant, and S. Previdi, “IP Fast ReRoute Using Not-via -addresses draft-ietf-rtgwg-ipfrr-notvia-addresses-05”, Network Working Group, Cisco Systems, Mar 2010.
- [SUR01] V. Suraev, “Global path recovery enhancement using Notify Reverse LSP”, IETF Internet Draft, Apr 2001.
- [UOF03] Iperf is free tools developed by University of Central Florida available online at <http://www.noc.ucf.edu/Tools/Iperf/> visited in Sep 2010.
- [ZAB06] A. Zaballos, and C. Seguí, “Analysis and simulation of IGP Routing Protocols”, University Ramon Llull (URL - La Salle Engineering) Barcelona, 2006.

Appendix A

Router Configuration

M10 Juniper router configuration for MPLS-TE

```
ottawalab@M10> show configuration
## Last commit: 2010-03-29 01:20:42 EDT by ottawalab
version 10.1R1.8;
system {
  host-name M10;
  time-zone America/Montreal;
  root-authentication {
    encrypted-password "$1$BZ71x/9p$6LfXhfThg.LB9Zv7Bmu0f0"; ## SECRET-DATA
  }
  login {
    user ottawalab {
      uid 2000;
      class superuser;
      authentication {
        encrypted-password "$1$ZtqeCKR0$b9ENZLHV6R3DDoHu2rcpd1"; ## SECRET-DATA
      }
    }
  }
  services {
    telnet;
  }
  syslog {
    user * {
      any emergency;
    }
  }
  file messages {
    any notice;
    authorization info;
  }
  ntp {
```

```
peer 10.10.10.1;
interfaces {
  ge-0/0/0 {
    description connected-to-PC1;
    unit 0 {
      family inet {
        address 10.1.1.1/30;
      }
      family mpls;
    }
    ge-0/1/0 {
      description connected-to-OMB;
      unit 0 {
        family inet {
          address 172.16.1.2/30;
        }
        family mpls;
      }
      fxp0 {
        unit 0 {
          family inet {
            address 192.168.63.11/24;
          }
        }
        lo0 {
          unit 0 {
            family inet {
              address 10.10.10.3/32;
            }
            protocols {
              rsvp {
                interface ge-0/0/0.0;
                interface fxp0.0 {
                  disable;
                }
              }
            }
          }
          mpls {
            label-switched-path "M10#2-to-M160" {
              to 10.10.10.2;
```

```
ldp-tunneling;
fast-reroute;
}
interface fxp0.0 {
disable;
}
interface ge-0/0/0.0;
interface ge-0/1/0.0;
}
ospf {
traffic-engineering;
area 0.0.0.0 {
interface ge-0/0/0.0;
interface ge-0/1/0.0;
interface lo0.0 {
passive;
}
}
ldp {
interface ge-0/1/0.0;
interface lo0.0;
```

M10#2 Router configuration for IP-FRR

```
## Last commit: 2010-07-12 14:21:24 EDT by ottawalab
version 10.2R1.8;
system {
  host-name M10;
  time-zone America/Montreal;
  root-authentication {
    encrypted-password "$1$BZ71x/9p$6LfXhfThg.LB9Zv7Bmu0f0"; ## SECRET-DATA
  }
  login {
    user ottawalab {
      uid 2000;
      class superuser;
      authentication {
```

```
        encrypted-password "$1$ZtqeCKR0$b9ENZLHV6R3DDoHu2rcpd1"; ## SECRET-DATA
    }
services {
    telnet;
}
syslog {
    user * {
        any emergency;
    }
    file messages {
        any notice;
        authorization info;
    }
}
ntp {
    peer 10.10.10.1;
}
interfaces {
    ge-0/0/0 {
        description connected-to-olive-1;
        unit 0 {
            family inet {
                address 10.1.1.1/30;
            }
            family mpls;
        }
    }
    ge-0/1/0 {
        description "connected-to-M10#1";
        unit 0 {
            family inet {
                address 10.1.1.9/30;
            }
            family mpls;
        }
    }
    at-0/2/0 {
        atm-options {
            pic-type atm1;
        }
    }
}
```

```
vpi 0 {
    maximum-vcs 256;
}
unit 0 {
    vci 0.128;
    family inet {
        address 10.0.89.6/32 {
            destination 10.0.89.5;
        }
    }
    family mpls;
}
ge-0/3/0 {
    description connected-to-OMB;
    unit 0 {
        family inet {
            address 172.16.1.2/30;
        }
        family mpls;
    }
}
fxp0 {
    unit 0 {
        family inet {
            address 192.168.63.11/24;
        }
    }
}
lo0 {
    unit 0 {
        family inet {
            address 10.10.10.3/32 {
                preferred;
            }
        }
    }
}
routing-options {
    traceoptions {
        file routing-log;
    }
}
forwarding-table {
    export PLB;
```

```
    }  
}  
protocols {  
    ospf {  
        traceoptions {  
            file ospf-log size 30k files 5;  
            flag spf;  
        }  
        traffic-engineering;  
        area 0.0.0.0 {  
            interface lo0.0 {  
                passive;  
            }  
            interface ge-0/3/0.0 {  
                no-eligible-backup;  
            }  
            interface ge-0/0/0.0 {  
                node-link-protection;  
            }  
            interface ge-0/1/0.0 {  
                metric 2;  
            }  
        }  
    }  
    ldp {  
        interface all;  
        interface fxp0.0 {  
            disable;  
        }  
    }  
    policy-options {  
        policy-statement PLB {  
            then {  
                load-balance per-packet;  
                accept;  
            }  
        }  
    }  
}
```

Appendix B

MPLS-TE Work Verification

Show 20 LSP Statues:

ottawalab@M10> show mpls lsp

Ingress LSP: 20 sessions

To	From	State	Rt P	ActivePath	LSPname
10.10.10.2	10.10.10.3	Up	0 *		M10#2-to-M160
10.10.10.2	10.10.10.3	Up	0 *		2
10.10.10.2	10.10.10.3	Up	0 *		3
10.10.10.2	10.10.10.3	Up	0 *		4
10.10.10.2	10.10.10.3	Up	0 *		5
10.10.10.2	10.10.10.3	Up	0 *		6
10.10.10.2	10.10.10.3	Up	0 *		7
10.10.10.2	10.10.10.3	Up	0 *		8
10.10.10.2	10.10.10.3	Up	0 *		9
10.10.10.2	10.10.10.3	Up	0 *		10
10.10.10.2	10.10.10.3	Up	0 *		11
10.10.10.2	10.10.10.3	Up	0 *		12
10.10.10.2	10.10.10.3	Up	0 *		13
10.10.10.2	10.10.10.3	Up	0 *		14
10.10.10.2	10.10.10.3	Up	0 *		15
10.10.10.2	10.10.10.3	Up	0 *		16
10.10.10.2	10.10.10.3	Up	0 *		17
10.10.10.2	10.10.10.3	Up	0 *		18
10.10.10.2	10.10.10.3	Up	0 *		19
10.10.10.2	10.10.10.3	Up	0 *		20

Total 20 displayed, Up 20, Down 0

Egress LSP: 0 sessions

Total 0 displayed, Up 0, Down 0

Transit LSP: 0 sessions

Total 0 displayed, Up 0, Down 0

Show MPLS LSP Fast ReRoute status:

ottawalab@M10> show mpls lsp detail

Ingress LSP: 1 sessions

10.10.10.2

From: 10.10.10.3, State: Up, ActiveRoute: 0, LSPname: M10#2-to-M160

ActivePath: (primary)

FastReroute desired

LSPTYPE: Static Configured

LoadBalance: Random

Encoding type: Packet, Switching type: Packet, GPID: IPv4

*Primary State: Up

Priorities: 7 0

SmartOptimizeTimer: 180

Computed ERO (S [L] denotes strict [loose] hops): (CSPF metric: 32)

10.1.1.2 S 10.1.1.6 S 10.1.1.14 S 10.1.1.18 S 10.0.1.2 S

Received RRO (ProtectionFlag 1=Available 2=InUse 4=B/W 8=Node 10=SoftPreempt 20=Node-ID):

10.1.1.2 10.1.1.6 10.1.1.14 10.1.1.18(flag=1) 10.0.1.2

Total 1 displayed, Up 1, Down 0

Egress LSP: 0 sessions

Total 0 displayed, Up 0, Down 0

Transit LSP: 0 sessions

Total 0 displayed, Up 0, Down 0

Calculation of Convergence time in case of MPLS-TE

ottawalab@M10> show mpls lsp extensive

Ingress LSP: 300 sessions

10.10.10.2

From: 10.10.10.3, State: Up, ActiveRoute: 0, LSPname: 68

ActivePath: (primary)

FastReroute desired

LSPTYPE: Static Configured

LoadBalance: Random

Encoding type: Packet, Switching type: Packet, GPID: IPv4

*Primary State: Up

Priorities: 7 0

SmartOptimizeTimer: 180

Computed ERO (S [L] denotes strict [loose] hops): (CSPF metric: 2)

10.1.1.2 S 10.1.1.22 S

Received RRO (ProtectionFlag 1=Available 2=InUse 4=B/W 8=Node 10=SoftPreempt 20=Node-ID):

10.1.1.2 10.1.1.22

85 Apr 1 04:06:44.301 Record Route: 10.1.1.2 10.1.1.22

84 Apr 1 04:06:44.301 Up

83 Apr 1 04:06:44.290 Originate make-before-break call

82 Apr 1 04:06:44.290 CSPF: computation result accepted 10.1.1.2 10.1.1.22

81 Apr 1 04:05:28.878 Record Route: 10.1.1.2 10.1.1.22

80 Apr 1 04:05:28.878 Up

79 Apr 1 04:05:28.466 Record Route: 10.1.1.2(flag=3) 10.1.1.22

78 Apr 1 04:05:27.189 CSPF: link down/deleted: 10.0.1.1(10.10.10.1:75)(10.10.10.1)-
10.0.1.2(10.10.10.2:0) (10.10.10.2)

77 Apr 1 04:05:27.088 Originate make-before-break call

76 Apr 1 04:05:27.088 CSPF: computation result accepted 10.1.1.2 10.1.1.22

75 Apr 1 04:05:26.850 10.1.1.2: Tunnel local repaired

74 Apr 1 03:49:43.475 Record Route: 10.1.1.2(flag=1) 10.0.1.2

73 Apr 1 03:39:54.926 Record Route: 10.1.1.2 10.0.1.2

72 Apr 1 03:39:54.926 Up

Convergence time = 04:05:28.878 - 04:05:26.850 = 2.028 sec

Appendix C

IP-FRR Work Verification

OSPF routing table having backup IPs

ottawalab@M10> show ospf route

Topology default Route Table:

Prefix	Path	Route	NH	Metric	NextHop	Nexthop
Type	Type	Type		Interface	Address/LSP	
10.10.10.1	Intra	Router	IP	2	ge-0/1/0.0	10.1.1.10
10.10.10.2	Intra	Router	IP	2	ge-0/0/0.0	10.1.1.2
		Bkup IP			ge-0/1/0.0	10.1.1.10
10.10.10.4	Intra	Router	IP	1	ge-0/0/0.0	10.1.1.2
10.10.10.20	Intra	Router	IP	3	ge-0/0/0.0	10.1.1.2
		Bkup IP			ge-0/1/0.0	10.1.1.10
10.10.10.21	Intra	Router	IP	1	ge-0/3/0.0	172.16.1.1
10.1.1.0/30	Intra	Network	IP	1	ge-0/0/0.0	
10.1.1.4/30	Intra	Network	IP	2	ge-0/0/0.0	10.1.1.2
		Bkup IP			ge-0/1/0.0	10.1.1.10
10.1.1.8/30	Intra	Network	IP	2	ge-0/1/0.0	
10.1.1.12/30	Intra	Network	IP	3	ge-0/0/0.0	10.1.1.2
					ge-0/1/0.0	10.1.1.10
10.2.2.1/32	Intra	Network	IP	2	ge-0/0/0.0	10.1.1.2
		Bkup IP			ge-0/1/0.0	10.1.1.10
10.2.2.2/32	Intra	Network	IP	2	ge-0/0/0.0	10.1.1.2
		Bkup IP			ge-0/1/0.0	10.1.1.10
10.2.2.3/32	Intra	Network	IP	2	ge-0/0/0.0	10.1.1.2
		Bkup IP			ge-0/1/0.0	10.1.1.10
10.2.2.4/32	Intra	Network	IP	2	ge-0/0/0.0	10.1.1.2
		Bkup IP			ge-0/1/0.0	10.1.1.10
10.10.10.1/32	Intra	Network	IP	2	ge-0/1/0.0	10.1.1.10
10.10.10.2/32	Intra	Network	IP	2	ge-0/0/0.0	10.1.1.2
		Bkup IP			ge-0/1/0.0	10.1.1.10

10.10.10.3/32	Intra Network	IP	0	lo0.0	
10.10.10.4/32	Intra Network	IP	1	ge-0/0/0.0	10.1.1.2
10.10.10.20/32	Intra Network	IP	13	ge-0/0/0.0	10.1.1.2
	Bkup IP			ge-0/1/0.0	10.1.1.10
10.10.10.21/32	Intra Network	IP	11	ge-0/3/0.0	172.16.1.1
172.16.1.0/30	Intra Network	IP	1	ge-0/3/0.0	
172.31.1.0/30	Intra Network	IP	3	ge-0/0/0.0	10.1.1.2
	Bkup IP			ge-0/1/0.0	10.1.1.10
192.168.1.0/24	Intra Network	IP	2	ge-0/3/0.0	172.16.1.1
192.168.2.0/24	Intra Network	IP	4	ge-0/0/0.0	10.1.1.2
	Bkup IP			ge-0/1/0.0	10.1.1.10
192.168.63.0/24	Intra Network	IP	3	ge-0/1/0.0	10.1.1.10

Eligible nodes to be covered by IP-FRR

ottawalab@M10> show ospf backup spf

Topology default results:

Area 0.0.0.0 results:

10.1.1.6

Self to Destination Metric: 2

Parent Node: 10.10.10.4

Primary next-hop: ge-0/0/0.0 via 10.1.1.2

Backup next-hop: ge-0/1/0.0 via 10.1.1.10

Backup Neighbor: 10.10.10.4

Neighbor to Destination Metric: 1, Neighbor to Self Metric: 1

Self to Neighbor Metric: 1, Backup preference: 0x0

Not eligible, Reason: Primary next-hop node fate sharing

Backup Neighbor: 10.10.10.1

Neighbor to Destination Metric: 2, Neighbor to Self Metric: 1

Self to Neighbor Metric: 2, Backup preference: 0x0

Eligible, Reason: Contributes backup next-hop

10.10.10.2

Self to Destination Metric: 2

Parent Node: 10.1.1.6

Primary next-hop: ge-0/0/0.0 via 10.1.1.2
Backup next-hop: ge-0/1/0.0 via 10.1.1.10
Backup Neighbor: 10.10.10.4
Neighbor to Destination Metric: 1, Neighbor to Self Metric: 1
Self to Neighbor Metric: 1, Backup preference: 0x0
Not eligible, Reason: Primary next-hop node fate sharing
Backup Neighbor: 10.10.10.1
Neighbor to Destination Metric: 1, Neighbor to Self Metric: 1
Self to Neighbor Metric: 2, Backup preference: 0x0
Eligible, Reason: Contributes backup next-hop

172.31.1.2
Self to Destination Metric: 3
Parent Node: 10.10.10.2
Primary next-hop: ge-0/0/0.0 via 10.1.1.2
Backup next-hop: ge-0/1/0.0 via 10.1.1.10
Backup Neighbor: 10.10.10.4
Neighbor to Destination Metric: 2, Neighbor to Self Metric: 1
Self to Neighbor Metric: 1, Backup preference: 0x0
Not eligible, Reason: Primary next-hop node fate sharing
Backup Neighbor: 10.10.10.1
Neighbor to Destination Metric: 2, Neighbor to Self Metric: 1
Self to Neighbor Metric: 2, Backup preference: 0x0
Eligible, Reason: Contributes backup next-hop

10.10.10.20
Self to Destination Metric: 3
Parent Node: 172.31.1.2
Primary next-hop: ge-0/0/0.0 via 10.1.1.2
Backup next-hop: ge-0/1/0.0 via 10.1.1.10
Backup Neighbor: 10.10.10.4
Neighbor to Destination Metric: 2, Neighbor to Self Metric: 1
Self to Neighbor Metric: 1, Backup preference: 0x0
Not eligible, Reason: Primary next-hop node fate sharing
Backup Neighbor: 10.10.10.1
Neighbor to Destination Metric: 2, Neighbor to Self Metric: 1

Self to Neighbor Metric: 2, Backup preference: 0x0

Eligible, Reason: Contributes backup next-hop

200 Prefixes Covered by IP-FRR

ottawalab@M10> show ospf backup coverage

Topology default coverage:

Node Coverage:

Area	Covered	Total	Percent
------	---------	-------	---------

Nodes	Nodes	Covered
-------	-------	---------

0.0.0.0	2	5	40.00%
---------	---	---	--------

Route Coverage:

Path Type	Covered	Total	Percent
-----------	---------	-------	---------

Routes	Routes	Covered
--------	--------	---------

Intra	200	209	95.69%
-------	-----	-----	--------

Inter	0	0	100.00%
-------	---	---	---------

Ext1	0	0	100.00%
------	---	---	---------

Ext2	0	0	100.00%
------	---	---	---------

All	200	209	95.69%
-----	-----	-----	--------

Calculation of IP-FRR convergence time

ottawalab@M10> show log ospf-log

Jul 5 16:11:09.562757 Finished flash processing for topology default

Jul 5 16:11:09.579394 RPD_OSPF_NBRDOWN: OSPF neighbor 10.1.1.2 (realm ospf-v2 ge-0/0/0.0 area 0.0.0.0) state changed from Full to Down due to KillNbr (**event reason: interface went down**)

Jul 5 16:11:09.629631 OSPF full SPF scheduled for topology default

Jul 5 16:11:09.629798 OSPF SPF scheduled for topology default in 0.200000s

Jul 5 16:11:09.831405 Starting full SPF for topology default

Jul 5 16:11:09.831715 SPF starting for area 0.0.0.0

Jul 5 16:11:09.832138 SPF finished for area 0.0.0.0

Jul 5 16:11:09.832257 SPF elapsed time 0.000524s

Jul 5 16:11:09.832316 Backup SPF's initialization starting for area 0.0.0.0

Jul 5 16:11:09.832420 Backup SPF's initialization finished for area 0.0.0.0 in 0.000098s

Jul 5 16:11:09.832535 Backup SPF's starting for area 0.0.0.0

Jul 5 16:11:09.832591 Backup SPF starting for area 0.0.0.0 root 10.10.10.1

Jul 5 16:11:09.832757 Backup SPF finished for area 0.0.0.0 root 10.10.10.1 in 0.000160s

Jul 5 16:11:09.832823 Backup SPF for 1 neighbors finished for area 0.0.0.0 in 0.000282s

Jul 5 16:11:09.832876 SPF result evaluation starting for area 0.0.0.0

Jul 5 16:11:09.832942 SPF result evaluation finished for area 0.0.0.0 in 0.000063s

Jul 5 16:11:09.832994 Backup SPF cleanup starting for area 0.0.0.0

Jul 5 16:11:09.833049 Backup SPF cleanup finished for area 0.0.0.0 in 0.000052s

Jul 5 16:11:09.833109 Backup SPF elapsed time 0.000793s

Jul 5 16:11:09.833163 SPF add stubs starting for area 0.0.0.0

Jul 5 16:11:09.833531 SPF add stubs finished for area 0.0.0.0

Jul 5 16:11:09.833595 Stub elapsed time 0.000429s

Jul 5 16:11:09.833653 SPF add inter area starting for area 0.0.0.0

Jul 5 16:11:09.833703 SPF add inter area Summary starting for area 0.0.0.0

Jul 5 16:11:09.833758 SPF add inter area Summary finished for area 0.0.0.0

Jul 5 16:11:09.833885 SPF add inter area ASBRSum starting for area 0.0.0.0

Jul 5 16:11:09.833939 SPF add inter area ASBRSum finished for area 0.0.0.0

Jul 5 16:11:09.833986 SPF add inter area finished for area 0.0.0.0

Jul 5 16:11:09.834037 Interarea elapsed time 0.000382s

Jul 5 16:11:09.834085 SPF add external starting

Jul 5 16:11:09.834130 SPF add external finished

Jul 5 16:11:09.834181 External elapsed time 0.000092s

Jul 5 16:11:09.834228 SPF add NSSA starting

Jul 5 16:11:09.834332 SPF add NSSA finished

Jul 5 16:11:09.834383 NSSA elapsed time 0.000150s

Jul 5 16:11:09.834431 SPF post spf cleanup starting

Jul 5 16:11:09.837986 SPF post spf cleanup finished

Jul 5 16:11:09.838084 Cleanup elapsed time 0.003645s

Jul 5 16:11:09.838138 Total elapsed time 0.006756s

Jul 5 16:11:09.838207 Finished full SPF for topology default

Jul 5 16:11:09.839742 Starting flash processing for topology default

Jul 5 16:11:09.839870 Finished flash processing for topology default

Jul 5 16:11:09.841095 Starting flash processing for topology default

Jul 5 16:11:09.841222 Finished flash processing for topology default

Total elapsed time (0.006756s) + Backup SPF elapsed time (0.000793s) = IP-FRR convergence time (7.549 ms)

Appendix D

Confidence Interval Calculations

Experiments data such as convergence time are measured by taking the mean of a succession of n tests, each of long enough time to ensure uncorrelated results. All tests are identical and independent from each other. The n independent results are represented by $X_1, X_2, X_3, \dots, X_{n-1}, X_n$.

$$\text{The Mean } \bar{X} = \frac{1}{n} \sum_{i=1}^n X_i \quad (\text{D.1})$$

However, the mean of the independent experiment tests $\bar{X}$ provide us with a single numerical value for the estimate of the expected value $E[X] = \mu$. In order to know how good is the estimate provided by $\bar{X}$ for the experiments results, it is necessary to compute the variance V_b^2 .

$$V_b^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 \quad (\text{D.2})$$

Small V_b^2 indicates that the results are tightly clustered around $\bar{X}$, and we can be confident that $\bar{X}$ is close to the $E[X]$. On the other hand, if V_b^2 is large, the results are widely dispersed about $\bar{X}$ and we can not be confident that $\bar{X}$ is close to the $E[X]$. Instead of seeking a single value to estimate the $E[X]$, we can specify the interval of values that is highly likely to contain the true value of the parameter. We begin by specifying some high probability, say $1-\alpha$. We then find the interval $[L(X), U(X)]$ such that the probability:

$$Pr[L(X) \leq \mu \leq U(X)] = 1 - \alpha \quad (\text{D.3})$$

The interval contains the true value of the parameters with probability $1 - \alpha$. Such an interval is $1 - \alpha \times 100\%$ confidence interval.

Using the standard deviation and t distributed table, the lower and upper limits of the 95% confidence interval can be calculated as follow:

$$\text{Lower limit } L(X) = \bar{X} - \frac{\sigma t_{[\frac{\alpha}{2}, n-1]}}{\sqrt{n}} \quad (\text{B.4})$$

$$\text{Upper Limit } U(X) = \bar{X} + \frac{\sigma t_{[\frac{\alpha}{2}, n-1]}}{\sqrt{n}} \quad (\text{B.5})$$

Where:

$$\alpha = 0.05$$

$$n = \text{number of tests}$$

$$\bar{X} = \text{sample average}$$

$$\sigma = \text{sample standard deviation} = \sqrt{V_x^2} = \sqrt{\frac{1}{n-1} \sum_i^n (X_i - \bar{X})^2}$$

The confidence interval means that 95% of the experiments results fall within the interval. Throughout this thesis, the confidence interval is computed based on three independent tests. From the table of t distribution, the $t_{[\alpha/2, \sigma]}$ is found to be 2.447. It is observed that more than 95% of the results are within the calculated confidence interval for each experiment. See Tables D.1, D.2, and D.3 for confidence interval calculation. The tables show the convergence times of MPLS-TE FRR and IP Fast ReRoute in Chapter 4 when experiments are carried out on the network topology in Figures 4.12, 4.22, and 4.25 with the increasing of some parameters. The average convergence time for the three tests are shown together with the calculated mean $\bar{X}$, the upper and lower values of interval $U(X)$ - $L(X)$. The confidence intervals are not shown on the figures because of the intervals are relatively small.

No. Of LSPs	Experiment tests averages			$\bar{X}$ msec	σ	U(X)	L(X)	Interval
	X_1 msec	X_2 msec	X_3 msec					
1	71.145	71.566	71.233	71.314	0.18131	71.51983	71.10949	0.41035
10	165.01	165.12	164.95	165.02	0.0704	165.1063	164.9470	0.043
20	241.54	241.77	241.44	241.58	0.13816	241.7396	241.4269	0.223
40	470.23	470.57	470.09	470.29	0.20155	470.5247	470.0685	0.077
80	840.23	840.65	840.17	840.35	0.21354	840.5916	840.1083	0.145
150	2100.9	2100.4	2100.6	2100.6	0.20548	2100.865	2100.400	0.066
300	4600.3	4600.5	4600.3	4600.3	0.09428	4600.473	4600.259	0.028

Table D.1: Confidence Interval Calculations in Multiple LSPs Scenario

No. Of Hops	Experiment tests averages			$\bar{X}$ msec	σ	U(X)	L(X)	Interval
	X_1 msec	X_2 msec	X_3 msec					
1	84.05	84.25	84.56	84.286	0.25697	84.57744	83.99588	0.58156
2	117.25	117.42	118.01	117.56	0.39801	118.0120	117.1112	0.90075
3	144.32	145.01	144.58	144.63	0.34847	145.0309	144.2423	0.78865
4	162.78	162.35	162.01	162.38	0.38588	162.8166	161.9433	0.87330
5	180.23	180.63	180.56	180.47	0.21362	180.7150	180.2316	0.48345

Table D.2: Confidence Interval Calculations in Multiple Hops Scenario

IGP Prefixes	Experiment tests averages			$\bar{X}$ msec	σ	U(X)	L(X)	Interval
	X_1 msec	X_2 msec	X_3 msec					
10	4.978	5.007	5.096	5.027	0.06149	5.062	4.992	0.070
20	5.785	5.853	5.789	5.809	0.038158	5.831	5.787	0.043
40	7.549	7.851	7.480	7.627	0.197318	7.389	7.516	0.223
50	8.263	8.302	8.17	8.245	0.067816	8.284	8.207	0.077
60	9.119	9.336	9.108	9.188	0.128578	9.261	9.116	0.145
80	10.893	10.795	10.79	10.826	0.058078	10.859	10.793	0.066
100	12.297	12.315	12.266	12.293	0.024786	12.370	12.279	0.028
120	14.436	14.289	14.45	14.392	0.089187	14.443	14.341	0.101
150	16.582	16.669	16.79	16.680	0.104462	16.739	16.621	0.118
200	21.088	20.783	20.749	20.873	0.186682	20.979	20.767	0.211

Table D.3: Confidence Interval Calculations in IP-FRR Scenario