

Universitatea “Politehnica” din București Facultatea de Electronică,
Telecomunicații și Tehnologia Informației

Optimizarea traficului de mare viteză între rețele de calculatoare

Proiect de diplomă

prezentat ca cerință parțială pentru obținerea titlului de Inginer în domeniul
Calculatoare și Tehnologia Informației programul de studii de licență
Ingineria Informației

Conducător științific

Conf. dr. ing. Ștefan STĂNCESCU

Absolvent

Gabriel GRĂDINARU-TAȘCĂU

2013

Universitatea "Politehnica" din București
Facultatea de Electronică, Telecomunicații și Tehnologia Informației
Departamentul Electronică Aplicată și Ingineria Informației

Aprobat Director de Departament :
Prof. Dr.Ing. Sever Pașca



TEMA PROIECTULUI DE DIPLOMĂ
a studentului
Grădinaru-Tașcău I. Gabriel, grupa 442A

1. Titlul temei: Optimizarea traficului de mare viteză între rețele de calculatoare
2. Contribuția practică, originală a studentului va consta în: comparație între mijloace de gestiune a traficului folosind simulatoare de rețea: NS (Network Simulator), Cisco Packet Tracer, OPNET, etc.; analiza unui mecanism de dirijare de trafic ATM și a unei arhitecturi ATM; optimizarea acestei arhitecturi și simularea funcționării acesteia; justificarea avantajelor soluției găsite; formularea de concluzii.
3. Proiectul se bazează pe cunoștințe dobândite în principal la următoarele discipline: Arhitecturi de Rețele și Internet, Rețele de Calculatoare, Sisteme de Comunicații
4. Realizarea practică/ proiectul rămân în proprietatea: Grădinaru-Tașcău Gabriel și UPB, Facultatea de Electronică, Telecomunicații și Tehnologia Informației
5. Proprietatea intelectuală asupra proiectului aparține: Grădinaru-Tașcău Gabriel și UPB, Facultatea de Electronică, Telecomunicații și Tehnologia Informației
6. Locul de desfășurare a activității: UPB, Facultatea de Electronică, Telecomunicații și Tehnologia Informației
7. Data eliberării temei: 15.11.2012

CONDUCĂTOR LUCRARE:

Conf. Dr. Ing. Ștefan Stăncescu



STUDENT:

Gabriel Grădinaru-Tașcău



Declarație de onestitate academică

Prin prezenta declar că lucrarea cu titlul “ Studiul și compararea unor simulatoare de rețea diverse și combaterea congestiei cu ajutorul unui algoritm cunoscut”, prezentată în cadrul Facultății de Electronică, Telecomunicații și Tehnologia Informației a Universității “Politehnica” din București ca cerință parțială pentru obținerea titlului de *Inginer* în domeniul *Calculatoare și Tehnologia Informației*, programul de studii *Ingineria Informației* este scrisă de mine și nu a mai fost prezentată niciodată la o facultate sau instituție de învățământ superior din țară sau străinătate.

Declar că toate sursele utilizate, inclusiv cele de pe Internet, sunt indicate în lucrare, ca referințe bibliografice. Fragmentele de text din alte surse, reproduse exact, chiar și în traducere proprie din altă limbă, sunt scrise între ghilimele și fac referință la sursă. Reformularea în cuvinte proprii a textelor scrise de către alți autori face referință la sursă. Înțeleg că plagiatul constituie infracțiune și se sancționează conform legilor în vigoare.

Declar că toate rezultatele simulărilor, experimentelor și măsurărilor pe care le prezint ca fiind făcute de mine, precum și metodele prin care au fost obținute, sunt reale și provin din respectivele simulări, experimente și măsurători. Înțeleg că falsificarea datelor și rezultatelor constituie fraudă și se sancționează conform regulamentelor în vigoare.

București, 15.12.1012

Absolvent Gabriel Grădinaru-Tașcău



Cuprins

Capitolul 1

1. Introducere.....	13
1.1 WAN(Wide Area Network).....	13

Capitolul 2

2. Tehnici de optimizare WAN.....	15
2.1 Introducere	15
2.2 Scăderea lăţimii de bandă.....	15
2.3 Repararea protocoalelor LAN care nu sunt potrivite pentru WAN.....	17
2.4 Folosirea compresiei pentru a reduce lăţimea de bandă.....	19
2.5 Suprapunerile redundante aduc fişierele mai aproape de utilizatori.....	20

Capitolul 3

3. Comparatie simulatoare de reţea.....	23
3.1 Introducere	24
3.2 Simulare şi emulare	24
3.3 Clasificarea simulatoarelor de reţea	25
3.4. Analiza unor simulatoare de reţea	26
3.4.1 NS2(Network Simulator version 2)	26
3.4.2 NS3(Network Simulator version 3)	28
3.4.3 OPNET(Optimized Network Engineering Tool)	28
3.4.4 OMNET++(Optical Micro-Networks Plus Plus)	29
3.4.5 Cisco Packet Tracer	30
3.4.6 JiST(Java in Simulation Time)	31

Capitolul 4

4. Tehnologii WAN.....	33
4.1 ATM (Asynchronous Transfer Mode)	33
4.1.1 Mecanismul de dirijare a traficului într-o reţea ATM.....	35
4.1.2 Segmentarea şi reasamblarea.....	36
4.2 MPLS (Multiprotocol Label Switching)	37
4.2.1 Concepte de bază MPLS.....	37
4.2.2 Potenţiale eşecuri ale resurselor.....	41
4.2.3 Obiective pentru supravieţuirea în caz de eşec.....	42
4.2.4 Avantaje MPS.....	43

Capitolul 5

5. Soluţii de recuperare MPLS.....	45
5.1 Redirijarea	46
5.2 Modelul lui Makam.....	46
5.3 Modelul lui Haskin (Reverse backup).....	47
5.4 Redirijarea rapida (Fast Rerouting).....	48
5.4.1 Redirijarea rapidă cu rezervă unu la unu	49
5.4.2 Redirijarea rapidă cu rezervă de facilitare (Fast reroute Facility Backup).....	51
5.5 Soluţia propusă de restabilire a traficului	52

Capitolul 6

6. Simulări.....	53
6.1 Simulatorul NS2.....	53
6.3 Scenariul simulat.....	55

Capitolul 7

7. Prezentarea rezultatelor	57
7.1 Timpul de întrerupere al serviciului folosind diverse metode de recuperare	58
7.2 Numărul de pachete pierdute folosind diverse metode de recuperare	59
7.3 Numărul de resurse rezervate folosind diverse metode de recuperare	60
7.4 Numărul de pachete pierdute în funcţie de rata de transfer	61

Concluzii.....	63
----------------	----

Bibliografie şi referinţe.....	65
--------------------------------	----

Anexe.....	67
------------	----

Lista acronimelor

WAN	Wide Area Network
LAN	Local Area Network
TCP	Transmission Control Protocol
ISP	Internet Service Provider
QoS	Quality of Service
IP	Internet Protocol
SSH	Secure Shell
DNS	Domain Name System
ATM	Asynchronous Transfer Mode
VCI	Virtual Channel Identifier
VPI	Virtual Path Identifier
CRC	Cyclic Redundancy Check
MPLS	Multiprotocol Label Switching
LSP	Label Switched Path
LER	Label Edge Router
LSR	Label Switching Router
FEC	Forward Equivalent Class
LDP	Label Distribution Protocol
VPN	Virtual Private Network
IGP	Interior Gateway Protocol

Lista figurilor

Figura 1.1: Relația dintre gazde și subrețea	13
Figura 2.1: Captură de pachete pentru afișarea unei pagini web	18
Figura 3.1: Simulatorul NS2	26
Figura 3.2: Simulatorul NS3	28
Figura 3.3: Simulatorul OPNET	29
Figura 3.4: Simulatorul OMNET++	30
Figura 3.5: Cisco Packet Tracer	31
Figura 4.1: Antetul celulei ATM	34
Figura 4.2: Căile virtuale sunt formate din canale virtuale	35
Figura 4.3: Cale comutată după etichetă (Label Switched Path - LSP)	38
Figura 4.4: Structura antetului MPLS	39
Figura 4.5: Funcționarea LDP (Label Discovery Protocol)	40
Figura 5.1: Modelul lui Makam	47
Figura 5.2: Modelul lui Haskin	48
Figura 5.3: Redirijarea rapidă cu rezervă unu la unu	50
Figura 5.4: Redirijarea rapidă cu rezervă unu la unu și unificare	51
Figura 5.5: Modelul de redirijare rapidă cu rezervă de facilitare	52
Figura 5.6: Soluția propusă de stabilire a LSP-urilor de rezervă	53
Figura 6.1: Arhitectura unui nod MPLS în NS2	53
Figura 6.2: Topologia utilizată	55
Figura 7.1: Timpul de întrerupere al serviciului metodelor de recuperare	58
Figura 7.2: Numărul de pachete pierdute folosind diverse metode de recuperare	59
Figura 7.3: Numărul de resurse rezervate folosind diverse metode de recuperare	60
Figura 7.4: Numărul de pachete pierdute în funcție de rata de transfer	61

Capitolul 1

1. Introducere

1.1 WAN (Wide Area Network)

Când vorbim despre trafic de mare viteză între rețele de calculatoare vorbim despre trafic în rețele de tip WAN (Wide Area Network).

O rețea WAN acoperă o arie geografică întinsă, o țară sau un întreg continent. Astfel de rețele comunică între ele prin Internet sau prin alte mijloace realizate de către companii de telefonie sau alți furnizori de servicii. Rețeaua WAN conține o colecție de mașini utilizate pentru a executa programele utilizatorilor sau gazde conectate între ele printr-o subrețea de comunicație. De cele mai multe ori, subrețeaua este compusă din linii de transmisie și elemente de comutare a mesajelor. Liniile de transmisie pot fi alcătuite din fire de cupru, fibră optică sau legături radio și au rolul de a transporta biții de date între gazde. Elementele de comutare sunt calculatoare specializate care conectează două sau mai multe linii de transmisie. Elementul de comutare are următorul rol: când primește date pe o anumită linie, trebuie să aleagă o nouă linie pentru a retransmite datele mai departe. De multe ori elementul de comutare se numește ruter (sau router).

În cele mai multe WAN-uri, rețeaua conține un număr mare de linii de transmisie, fiecare dintre ele legând o pereche de rutere. Dacă două rutere nu sunt conectate direct prin același cablu, ele pot comunica, dar indirect, prin intermediul altor rutere. Când un pachet este transmis de la un ruter la altul prin intermediul unuia sau mai multor rutere, pachetul este primit de fiecare ruter intermediar, păstrat până când linia de ieșire dorită este liberă și apoi este retransmis. O subrețea care funcționează pe acest principiu se numește subrețea cu comutare de pachete. Aproape toate rețelele WAN au subrețele cu comutare de pachete. Atunci când pachetele sunt de dimensiune mică și au toate aceeași mărime, ele sunt adesea numite celule (în cazul ATM).

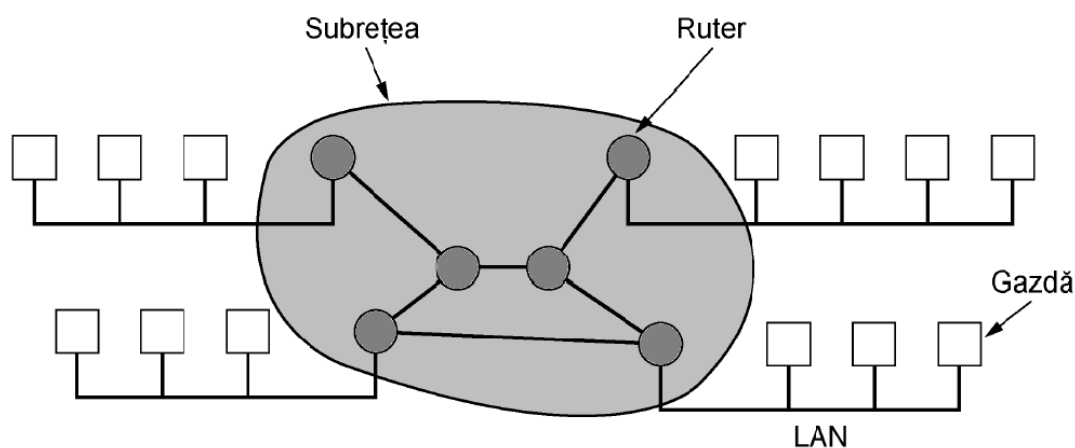


Figura 1.1: Relația dintre gazde și subrețea[1]

În general, atunci când un proces al unei gazde vrea să transmită un mesaj către un proces de pe o altă gazdă, gazda care transmite va împărți mesajul în pachete, fiecare dintre ele reținându-și numărul de ordine din secvență. Aceste pachete sunt apoi transmise prin rețea unul câte unul într-o succesiune rapidă și când ajung la gazda receptoare, aceasta le depozitează după care sunt reasamblate în mesajul inițial și furnizate procesului receptor.

Există și WAN-urile care nu sunt cu comutare de pachete. Poate exista un WAN format dintr-un sistem de sateliți. Fiecare ruter are o antenă prin care poate trimite și primi date. Toate ruterele pot asculta ieșirea altui satelit, iar în unele cazuri pot să asculte chiar și transmisia celorlalte rutere către satelit. Câteodată, ruterele sunt conectate la o rețea punct-la-punct și numai unele dintre ele dețin antene de satelit. De obicei, rețelele satelit sunt rețele cu difuzare și sunt foarte utile când proprietatea de difuzare este importantă.

Capitolul 2

2. Tehnici de optimizare WAN

2.1 Introducere

Optimizarea WAN este o colecție de tehnici destinate a crește eficiența transferurilor de date peste rețele WAN. Cele mai uzuale criterii de măsură a eficienței transferurilor TCP sunt rata de transfer, lățimea de bandă, latența, optimizarea protocoalelor și congestia, care se manifestă prin pierderea de pachete. De asemenea, și WAN-ul se poate clasifica după distanța dintre nodurile terminale și cantitatea de date transferate.

Companiile comerciale folosesc două topologii de rețele, prima este de tipul sediu și agenție iar cea de-a doua este de tipul centru de date către centru de date. În general, legăturile agențiilor cu WAN-ul sunt la distanța mică, folosesc mai puțină lățime de bandă, suportă mai multe conexiuni simultan, suportă conexiuni mai mici și de mai scurtă durată și sunt traversate de o mare varietate de protocoale. Sunt folosite pentru aplicații de afaceri cum ar fi poșta electronică, sisteme de gestionare a conținutului, aplicații pentru baze de date și acces către internet. În comparație, legăturile centrelor de date cu WAN au tendința să folosească mai multă lățime de bandă, sunt la distanțe mai mari, și au mai puține conexiuni dar mai mari (flux-uri de 100Mb/s până la 1Gb/s) și de durată mare. Traficul între centre de date poate include replicare de baze de date, back up, migrări de date, virtualizare și alte flux-uri specifice bazelor de date.

Optimizarea WAN a fost subiectul de cercetare academică intensă de când a apărut WAN-ul. La început s-au concentrat eforturile pe optimizarea conexiunilor WAN-ului cu agențiile, dar mai recent, creșterea rapidă a datelor digitale și nevoia de a le stoca și proteja, au stat la baza nevoii de optimizare a legăturilor centrelor de date cu WAN-ul.

2.2 Scăderea lățimii de bandă

Tratarea simptomelor și nu a cauzelor este o soluție comună, dar de cele mai multe ori greșită pentru multe probleme, inclusiv utilizarea lățimii de bandă. Rezultatele pot părea pozitive la început dar de obicei, de scurtă durată, neprevizibile și nu scalează bine. Scăderea lățimii de bandă reflectă această abordare de tratare a simptomelor (transferuri încete și timpi de răspuns excesivi) decât cauzele (puțină lățime de bandă disponibilă,

aplicații neautorizate, legături către resurse necorespunzătoare etc.). Totuși, această tehnică încă se folosește și încă destul de mult.

Scăderea lățimii de bandă asigură faptul că dispozitive care folosesc multă lățime de bandă, precum routerele sau gateway-urile, limitează cantitățile de date transmise și recepționate pe anumite perioade de timp. Scăderea lățimii de bandă limitează congestia rețelei și reduce instabilitatea server-ului rezultată din saturația rețelei. Pentru furnizorul de internet (Internet Service Provider - ISP), scăderea lățimii de bandă restricționează viteza utilizatorului pentru anumite aplicații sau în timpul perioadelor de folosire la maxim. Dar fără a înțelege cauzele creșterii traficului care determină scăderea lățimii de bandă, este doar o chestiune de timp până când această creștere va reveni și iar se va scădea lățimea de bandă.

Serverele care conțin date funcționează pe un principiu simplu de cerere și ofertă: clienții fac cereri și serverul le răspunde. Dar serverele care furnizează date în internet sunt predispuse supraîncărcării în timpul orelor de vârf și trebuie să prelucreze foarte multe cereri. Aceste perioade de încărcare maximă produc congestie sau îngreunează conexiunea care poate cauza instabilitatea serverului și o eventuală cădere a sistemului, rezultând un timp de inactivitate. Scăderea lățimii de bandă este folosită ca o metodă preventivă de a controla nivelul de răspuns al serverului la creșterea cererilor din partea clienților în timpul orelor de vârf.

În februarie 2008, membrii Comisiei Federale de Comunicații (Federal Communications Commission) au anunțat că se consideră stabilirea unor regulamente pentru a descuraja furnizorii de internet să micșoreze lățimea de bandă de la anumite site-uri și servicii care sunt mari consumatoare de lățime de bandă. Organizațiile pot (și ar trebui) să micșoreze lățimea de bandă sau filtre firewall pentru a limita sau a bloca traficul care violează politica de uz acceptabil (Acceptable Use Policy). Dar în alte conditii, micșorarea lățimii de bandă este aplicată cel mai bine sub forma clasei serviciului (Class of Service - CoS) sau calitatea serviciului (Quality of Service - QoS) aplicate anumitor tipuri de trafic. CoS și QoS reprezintă scheme de clasificare a traficului rețelei care acordă prioritate traficului dependent de timp sau traficului critic pentru misiune (mission-critical) decât să limiteze un anumit tip de trafic. Mulți experți recomandă ca pentru protocoale neautorizate sau nedorite să le fie micșorata lățimea de bandă la nivele foarte mici (sub 10Kb/s) decât să fie blocate complet, pentru a oferi administratorilor de rețea șansa de a le înlătura și a rezolva cu utilizatorii sau programele implicate. De exemplu, limitând lățimea de bandă disponibilă pentru protocoale peer-to-peer, precum BitTorrent(folosit pentru video și alte descărcări media personale), la 5Kb/s sau 10Kb/s,

administratorii ar putea avea timp să identifice stațiile sau serverele care sunt puncte terminale pentru activități peer-to-peer și să identifice indivizii care le utilizează. Pot să consilieze sau să disciplineze utilizatorii care sunt politicile de uz acceptabil.

2.3 Repararea protocoalelor LAN care sunt nepotrivite pentru WAN

Multe protocoale care mențin funcții vitale interne și externe ale unei afaceri funcționează cu scop limitat sau scalabilitate limitată. Acestea sunt, de cele mai multe ori, protocoale LAN învechite care pot impune o mare încetinire asupra rețelei WAN. De exemplu, sistemul de fișiere comune pe internet (Common Internet Files System - CIFS) crește exponențial în timp ce este transportat printr-un WAN. De asemenea, fluxurile de date în timp real și protocoalele de voce introduc nevoia pentru modelarea traficului. Pentru aceasta, traficul în rețea este controlat pentru a optimiza performanța, a satisface garanțiile de performanță, sau pentru a crește lățimea de bandă utilizabilă, de obicei prin întârzierea pachetelor care pot fi categorizate ca relativ insensibile la întârzieri sau care îndeplinesc anumite criterii de clasificare care marchează acest trafic cu prioritate scăzută. Modelarea traficului este folosită pentru a oferi utilizatorilor servicii voce și video satisfăcătoare, în timp ce se menține rețeaua în mod obișnuit pentru celelalte protocoale.

Urmărind aceeași idee, criptarea și aplicații de accelerare a protocoalelor de securitate devin mari consumatoare de resurse dar sunt poveri necesare, în special când traficul sensibil trebuie să traverseze prin internet. Până și cel mai folosit protocol din internet, HTTP-ul, poate fi descris ca având comunicații frecvente și constante și în același timp având perioade în care sunt cereri de resurse excesive.

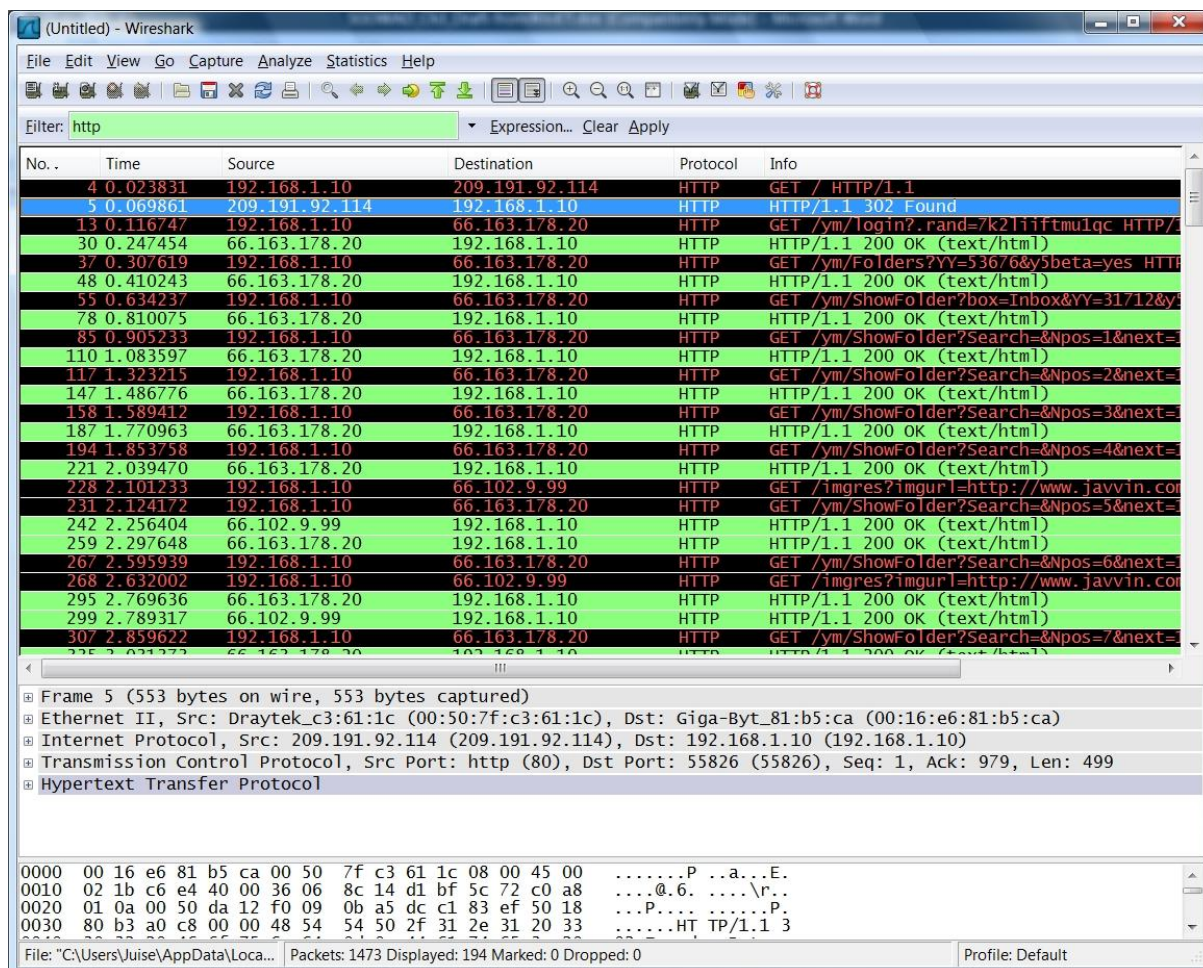


Figura 2.1: Captură de pachete pentru afișarea unei pagini web

Repararea acestor aspecte ale mediului rețelei devine o propunere de soluționare a traficului care ia în considerare nu doar aplicațiile în sine dar și programarea aplicațiilor în general. Știind cum funcționează o aplicație, formatul și parametrii protocoalelor sale și comportamentul în timpul rulării este esențial pentru înțelegerea integrării sale cu alte aplicații, servicii și protocoale din rețea. Nu este doar o propunere de reparare care presupune rezolvarea componentelor individuale, dar care cere folosirea celor mai bune practici pentru comunicații WAN eficiente: transmite și recepționează rar, în cantități mari, și sub forma unor tranzacții complete pe cât posibil.

Formatul TCP a fost proiectat inițial să funcționeze în mod fiabil peste medii de transmisie nesigure, fără a depinde de ratele de transmisie, întârzieri, coruperea protocolului, date duplicate și reordonarea segmentelor. Din cauza aceasta, TCP este întradevăr un mecanism robust și fiabil pentru transportul aplicațiilor, serviciilor și protocoalelor. Dar acest avantaj de design scoate la iveală dezavantajul TCP când este folosit peste o rețea modernă, medii de transmisie de mare viteză care întrec cu mult condițiile sub care TCP a fost proiectat să funcționeze.

Protocolul IP trimite pachete fără să se asigure dacă au ajuns la destinație; TCP menține conexiunile în desfășurare de la un capăt la altul, de la începutul transmisiei până la terminarea acesteia, are nevoie și de confirmări periodice că au fost transmise datele. Datele care nu au fost confirmate determină un algoritm de back-off exponențial după care se reîncearcă transmisia până când este recepționată și confirmată, sau trece o anumită perioadă de timp și semnalizează căderea conexiunii. Există un termen, sliding TCP window sizes, care semnifică numărul de pachete care poate fi trimis până când se recepționează confirmarea primirii pachetelor și care influențează în mod direct performanța când valori mai mari determină rate de transfer mai mari (dar și potențiale întârzieri mai mari). TCP folosește un algoritm bine definit de început lent care inițiază comunicația cu window size-ul foarte mic, după care scalează acest număr la proporții optime în timp ce conexiuni sunt stabilite și menținute active. Fiecare dintre acestea și alte proceduri asemănătoare ale stivei TCP/IP introduc întârzieri în rețea adresate în soluții de optimizare WAN prin tehnici de optimizare a conexiunii.

2.4 Folosirea compresiei pentru a reduce lățimea de bandă

Ca rezultat al abilității lor de a reduce volumul de trafic, mai multe scheme de compresie destinate îmbunătățirii performanței au fost o parte integrală din comunicațiile între rețele încă de pe vremea protocolului de transfer de fișiere ZMODEM introdus în 1986. Chiar și după 20 și de ani compresia apare în mod regulat în rețele moderne cu protocoale moderne precum Secure Shell (SSH) și compresia octeților din fluxurile de date în produse de optimizare WAN contemporane.

În termeni de telecomunicații, compresia lățimii de bandă înseamnă: o reducere a nevoii de lățime de bandă necesară pentru a transmite o anumită cantitate de date într-un anumit timp; sau o reducere în timpul necesar transmiterii unei anumite cantități de date cu o anumită lățime de bandă disponibilă. Aceasta implică o reducere în lățimea de bandă normală (sau timpul) pentru semnalele purtătoare de informații fără a reduce conținutul informației, datorită folosirii adecvate a tehnicilor de compresie de date. Acestea sunt bune și frumoase într-un mediu WAN, dar pot fi ineficiente fără acces către hardware de accelerare a compresiei pentru a obține compresie și decompresie maximă cu întârzieri minime (cu toate că software-ul este rapid, hardware-ul este, de obicei, cu câteva ordine de mărime mai rapid). În general, dispozitivele de accelerare WAN includ simboluri și dicționare pentru a reduce volumul datelor chiar mai mult decât poate compresia.

Redundanța repetitivă, care descrie modurile în care șabloanele și elementele tind să se repete în fluxurile de trafic regulat între perechi de transmițători și receptori, rămâne motivul esențial pentru care schemele de compresie sunt valabile și eficiente. Dintr-o perspectivă analitică corectă, multe date care traversează rețeaua includ repetiție inutilă, care face risipă de biți, timp și lățime de bandă necesară pentru transportul lor. Compresia datelor prin toate mijloacele posibile ajută la restaurarea echilibrului rețelei și este doar una dintre mai multe metode împotriva întârzierii nedorite.

Mai multe tehnici de compresie sunt potrivite pentru mediul rețelilor, dar toate se străduiesc să reducă lățimea de bandă consumată în timp ce traversează WAN-ul. Tehnici de compresie ale antetului și încărcăturii folosesc algoritmi de asemănare a șabloanelor pentru a identifica serii de octeți scurte, frecvente și repetitive din rețea care sunt înlocuite cu segmente mai scurte de cod pentru a reduce dimensiunea finală de transmisie. Algoritmi simplificați identifică șabloane de octeți din pachetele individuale pe când formele sofisticate de compresie pot analiza șabloane pe mai multe pachete și fluxuri de date.

Orice câștig oferit de strategiile de compresie trebuie să varieze conform tipului de trafic din rețeaua WAN. Arhivele comprimate de date (fișiere ZIP sau tar) nu pot fi reduse prea mult folosind scheme de compresie specifice rețelilor, dar aplicarea compresiei asupra diverselor fluxuri de date poate să mărească lățimea de bandă efectivă din WAN. Protocoalele de voce au parte de îmbunătățiri semnificative de pe urma compresiei antetului UDP în combinație cu alte tehnici precum combinarea pachetelor, descrisă în secțiunea următoare.

2.5 Suprapunerile redundante aduc fișierele mai aproape de utilizatori

Redundanța nu este întotdeauna un lucru rău pentru performanța rețelilor sau pentru ratele de transfer. Există multe aplicații unde multiplicitatea poate îmbunătăți performanța. De fapt mediile WAN au mari beneficii ca urmare a suprapunerilor redundante de servicii de fișiere care aduc datele și fișierele mai aproape de utilizatori.

Așa cum ne-a aratat sistemul DNS, o rețea poate funcționa eficient în moduri consolidate și descentralizate în același timp. Un bloc uniform de informație poate fi consolidat de la mai multe surse și distribuit printr-un sistem complet descentralizat pentru transmitere oriunde există o rețea. Consolidarea datelor agregate într-un singur centru de date creează beneficii pentru acei utilizatori care sunt apropiați de centrul de date, dar poate crea dificultăți de accesare pentru alți utilizatori care sunt mai departați, în special cei care trebuie să folosească conexiuni WAN înguste sau scumpe pentru a se conecta cu acel centru de date.

Apelând la o autoritate centrală sau la o sursă de informații pentru o companie dispersată pe tot globul pot apărea probleme pentru utilizatorii la distanță sau care se conectează din alte rețele, deci este mai bine să reproducă informația în locuri de unde pot fi accesate rapid de către utilizatori, oriunde s-ar afla aceștia. Bazele de date DNS, cu versiunile sale cu autoritate mai mare sau mai mică, cu principiile sale master și slave, împreună cu folosirea depozitelor cache a activității recente creează un model pentru distribuirea și dispersarea informației care funcționează destul de bine încât mențin internetul global. În mod similar, suprapunerile redundante caută să păstreze fișierele pe care utilizatorii le accesează nu mai departe de una sau doua hop-uri WAN distantă de mașinile lor, indiferent unde s-ar afla și la orice moment.

Astăzi, multe companii se străduiesc cu consolidarea și multiplicarea serverelor. Numeroase servere locale trebuie să găzduiască o varietate de servicii și protocoale și în ciuda numărului lor ridicat, acestea nu sunt întotdeauna în locul potrivit și la momentul potrivit. În schimb, multe companii preferă să instaleze câteva servere cheie în locații strategice geografice și teritoriale pentru a acomoda mai ușor utilizatori și echipe „în deplasare”. Această abordare permite unui grup de servere sincronizate cu grijă să furnizeze accesibilitate, lățime de bandă și securitate comparabilă oricui și oriunde. Reproducerea pe multiple servere fizice face această abordare posibilă, în timp ce virtualizare face posibilă rularea serviciilor în spații virtuale separate și face întregul proces mai practic ușurează întreținerea și monitorizarea.

Resursele rețelei și utilizatorii sunt de multe ori dispersați pe mai multe agenții, în locațiile clienților și pe mai mulți furnizori de internet care se întind pe regiuni multiple. O strategie corectă de consolidare a serverelor centralizează, în mod necesar, infrastructura și reduce numărul de servere pentru a reduce costul și a îmbunătăți organizarea. Din păcate, acest lucru creează nevoia de a asigura conectivitate neîntreruptă între locații la distanță pe legături cu WAN-ul și livrarea datelor eșuează atunci când aceste legături cedează. Înseamnă că trebuie să existe un mecanism sau mai multe care să prelucreze traficul atunci când numărul utilizatorilor conectați la distanță crește.

În timpul orelor de deschidere, multe afaceri suferă de o creștere de trafic inițial care este format, în mare parte, din utilizatorii care se înregistrează simultan pe unul sau mai multe servere. Autentificarea și accesarea serviciilor DNS este problematică în timpul acestei rutine în care toată lumea apare și se înregistrează în același timp, deci trebuie să existe un mod de a optimiza și prioritiza traficul astfel încât timpii de înregistrare să fie minimi.

O soluție consolidată de optimizare WAN ajută la obținerea celei mai bune utilizări a arhitecturii rețelei deoarece oferă o posibilitate de a profita la maxim de lățimea de bandă pe care o are o legătură WAN, în timp ce se optimizează traficul cu prioritate astfel încât, chiar și atunci când are loc o congestie, acest trafic să ajungă la destinație. Prin urmare, utilizatorul de dimineata ar putea aștepta chiar și un minut să descarce un e-mail, dar tot o să obțină un răspuns rapid atunci când o să furnizeze un cont de utilizator și o parolă pentru a se înregistra.

Capitolul 3

3. Comparație simulatoare de rețea

3.1 Introducere

În domeniul rețelelor, stabilirea unui scenariu de utilizare în timp real este foarte dificil. Pentru a reproduce fizic o rețea cu echipamente reale necesită un cost foarte mare, timp și bani. Așadar implementarea unei rețele întregi nu este un lucru ușor de făcut. Un simulator de rețea ajută administratorul să verifice dacă rețeaua este capabilă să lucreze în timp real. Prin urmare, timpul și costul testării funcționalității rețelei este redus și implementarea este facilitată. În următoarele pagini vom prezenta principalele caracteristici ale unor simulatoare de rețea și vom analiza avantajele și dezavantajele lor. Această analiză este utilă pentru a ajuta un proiectant sau administrator de rețea în alegerea tipului de rețea potrivit nevoilor sale.

Simularea este una dintre cele mai importante tehnologii din era modernă. Simulările computerizate pot modela obiecte reale sau ipotetice pentru a fi studiate. De asemenea, o rețea de calculatoare poate fi simulată pe un computer.

Un simulator de rețea este o tehnică de implementare a rețelei pe un computer care permite cercetătorilor să testeze scenarii care sunt dificile sau costisitoare de implementat în lumea reală. Este util în special pentru a testa protocoale noi sau pentru a modifica protocoalele existente într-un mediu de control și reproductibil. Se pot crea diferite tipologii folosind tipuri variate de noduri (stații de lucru, switch-uri, routere, unități mobile, etc). Simulatoarele de rețea sunt folosite de persoane din diverse arii, cum ar fi cercetări academice, software development și controlul calității pentru a proiecta, verifica și analiza performanța diverselor protocoale de rețea. În general, un simulator de rețea este format dintr-o varietate de tehnologii și protocoale și ajută utilizatorii să construiască rețele complexe din blocuri de bază cum ar fi clustere de noduri și legături. De menționat că simulatoarele de rețea nu sunt perfecte. Ele nu pot modela perfect toate detaliile rețelelor. Totuși, dacă sunt corect modelate pot oferi cercetătorului o perspectivă folositoare asupra rețelei testate și cum modificările pot afecta funcționalitatea acesteia.

3.2 Simulare și emulare

În domeniul cercetării rețelelor de calculatoare și comunicații, simularea este o tehnică utilă deoarece comportamentul unei rețele poate fi modelat prin calcularea interacțiunii dintre componentele rețelei folosind formule matematice. O rețea mai poate fi modelată prin captarea și reproducerea de observații experimentale din rețele reale aflate în producție. După ce obținem datele dintr-un experiment de simulare, comportamentul rețelei și protocoalele suportate pot fi observate și analizate într-o serie de teste fără a mai avea nevoie de rețeaua reală. Pot fi modificate diverse atribute ale mediului într-un mod controlat pentru a vedea cum se comportă rețeaua în condiții diferite sau cu diverse combinații de parametri.

Emularea unei rețele înseamnă că o anumită rețea este simulată pentru a evalua performanța ei sau pentru a prezice impactul unor posibile modificări sau optimizări. Diferența majoră între simulare și emulare sunt sistemele terminale, precum computerele, pot fi atașate emulatorului și se vor comporta ca și cum s-ar afla într-o rețea reală. Scopul unui emulator de rețea este să emuleze o rețea care interconectează stații terminale, dar nu emulează și aceste stații terminale. Anumite aplicații de emulare includ NS2, care este în sine un simulator dar poate fi utilizat și ca un emulator cu funcționalitate limitată. În contrast, un emulator de rețea veritabil este WANSim, care emulează o rețea WAN și care utilizează o parte din funcționalitatea Linux.

3.3 Clasificarea simulatoarelor de rețea

Există diverse simulatoare de rețea cu diverse caracteristici cum ar fi: NS2, NS3, OPNET, NetSim, OMNet++, J-Sim, QualNet, Cisco Packet Tracer. Acestea se pot clasifica pe baza unor criterii precum prețul și complexitatea.

Există simulatoare comerciale care nu pot fi folosite gratuit, utilizatorii trebuie să plătească pentru a achiziționa o licență pentru întreg software-ul sau doar pentru anumite pachete specifice. Exemple: OPNET, QualNet. Avantajul acestor simulatoare este documentația completă și actualizată care este menținută de către personal specializat dintr-o companie. În schimb, avantajul simulatoarelor gratuite și open source este faptul că orice persoană sau organizație poate contribui la dezvoltare și la găsirea bug-urilor. De asemenea interfața poate fi îmbunătățită și poate reflecta dezvoltări tehnologice recente mai rapid decât un simulator comercial. Exemple: NS2, NS3, OMNet++, J-Sim.

Simulatoarele de rețea existente sunt foarte variate începând de la unele foarte simple până la cele foarte complexe. În mod minimal, un simulator ar trebui să permită

utilizatorilor să reprezinte o topologie de rețea, să definească scenariile, să specifice nodurile rețelei, legăturile dintre noduri și traficul dintre noduri. Sistemele mai complicate pot permite utilizatorului să specifice toate aspectele protocoalelor folosite pentru a procesa traficul rețelei. De asemenea, aplicațiile grafice permit vizualizarea funcționalității mediului simulat. Unele pot fi bazate pe text și furnizează o interfață mai puțin vizuală sau intuitivă, dar pot permite forme mai avansate de personalizare. Altele pot fi orientate pe programare și pot oferi un framework care permite utilizatorilor crearea unei aplicații care să simuleze mediul de testare.

Cel mai utilizat și popular simulator de rețea este NS2. Acesta și-a stabilizat poziția ca fiind standardul simulatoarelor de rețea deoarece are foarte multe librării și foarte mulți au participat la dezvoltarea sa. Organizații non-profit au contribuit cu componente în librăriile sale și s-a demonstrat că modul de dezvoltare al NS2 a fost un succes. Cu toate acestea, din cauza unor limitări de design, a fost nevoie să se dezvolte succesorul său, NS3, care pune accentul pe documentație și niște persoane foarte specializate s-au oferit voluntare pentru a se ocupa de diferite componente. Totuși, NS3 nu este doar o versiune actualizată a NS2. Au fost reimplementate mai multe mecanisme bazate pe experiențe reușite sau mai puțin reușite ale NS2.

3.4 Analiza unor simulatoare de rețea

3.4.1 NS2(Network Simulator version 2)

NS2 este un simulator bazat pe evenimente discrete orientat către cercetare. Oferă suport pentru TCP, rutare și protocoale multicast peste toate tipurile de rețele. A fost dezvoltat în 1995 în proiectul VINT (Virtual Inter Network Testbed).

Simulările NS2 sunt compuse din cod C++ care modelează comportamentul nodurilor și scripturi oTcl care controlează simularea și alte aspecte cum ar fi topologia rețelei. Această alegere de design a fost făcută pentru a evita recompilări inutile dacă se produc schimbări în configurația simulării. În anul 1996, când a fost lansată prima versiune NS2, această alegere era rezonabilă deoarece recompilările frecvente C++ erau mari consumatoare de timp dar pentru standardele actuale, acest design al NS2 îl împiedica să fie eficient în simularea rețelelor la o scară mare. Pentru a anima simulările este folosit Nam (Network Animator), un tool scris în Tcl/TK, care suportă animarea la

nivel de pachet, topologia rețelei și alte tool-uri de inspecție a pachetelor. De obicei, datele de intrare pentru Nam sunt provenite din simulări NS2 sau pot proveni din măsurători dintr-o rețea reală, de exemplu din tcdump sau wireshark.

Cea mai recentă versiune de NS2 este NS-2.35 și a fost lansată în Noiembrie 2011. Această versiune are peste 300.000 linii de cod și foarte probabil încă pe atât cod contribuit de utilizatori care nu a fost integrat în versiunea principală. NS2 poate fi rulat pe următoarele platforme: Linux, FreeBSD, Solaris, Mac OS X și Windows.

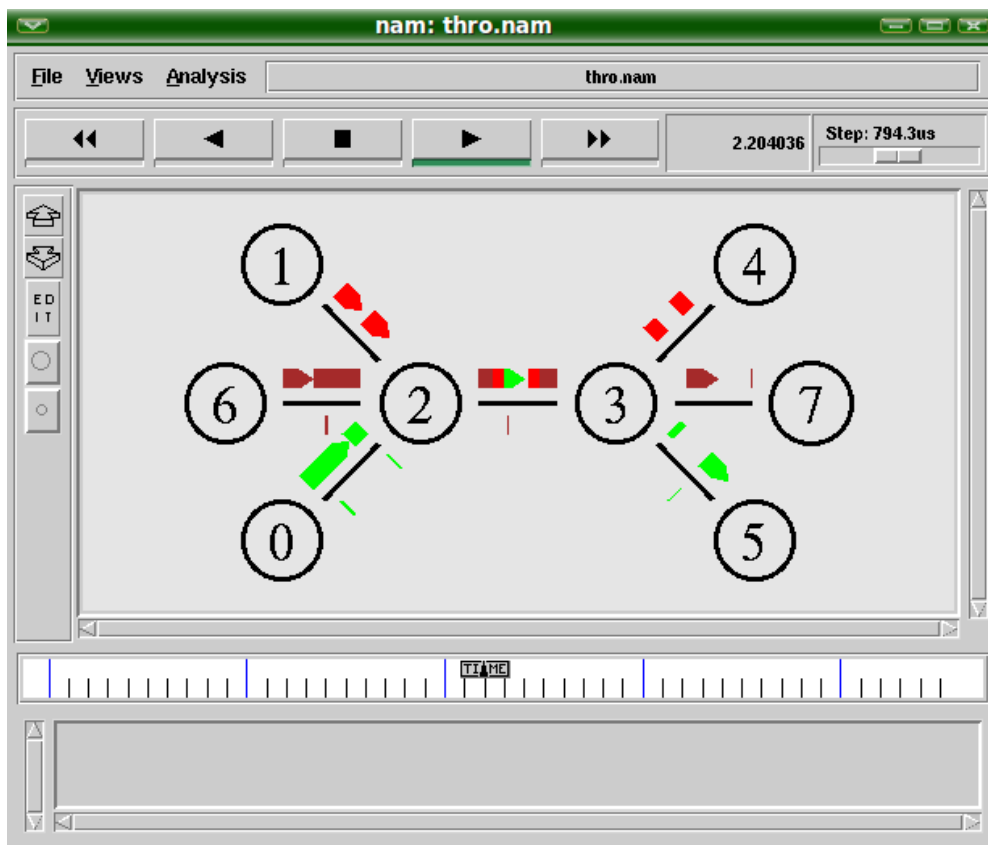


Figura 3.1: Simulatorul NS2

3.4.2 NS3 (Network Simulator version 3)

La fel ca predecesorul său, NS3 se bazează pe C++ pentru implementarea nodurilor simulării. Totuși, NS3 nu mai folosește scripturi oTcl pentru a controla simularea, astfel abandonând probleme introduse de combinația celor două limbaje, C++ și oTcl din NS2. În schimb, simulările din NS3 pot fi implementate în limbaj pur C++, în timp ce unele părți ale simulării pot fi realizate opțional folosind limbajul Python. În plus, NS3 integrează concepte arhitecturale și cod din GTNetS, un simulator cu capacități bune de scalabilitate. Aceste decizii de design au fost făcute cu prețul compatibilității, adică modelările din NS2 trebuie să fie portate manual.

Pe lângă îmbunătățirea performanței, au mai fost adăugate caracteristici noi precum socket-uri Berkley sau thread-uri POSIX. Proiectul NS3 a început în anul 2006, este gratuit, open-source și destinat pentru cercetare și dezvoltare.

Comparat cu NS2, avem următoarele diferențe majore între cele două simulatoare:

- Atenție sporită către realism: entitățile de protocol sunt proiectate astfel încât să fie mai apropiate de computerele reale;
- Integrare Software: suport încorporat pentru mai mult software de rețelistică open source și astfel se reduce nevoia de a rescrie modele pentru simulare;
- Suport pentru virtualizare: sunt folosite mașini virtuale mai puțin performante pentru a crea medii de testare potrivite.
- Arhitectura de analiza a datelor: NS3 construiește un framework care să gestioneze datele și statisticile astfel încât să poată fi prelucrate datele simulării.

Există totuși câteva limitări a simulatorului de rețea NS3, la baza acestora stă nevoia de participare a comunității de cercetare. O limitare generală a simulărilor este credibilitatea, care ar putea fi îmbunătățită și la NS3, lucru realizabil prin folosirea NS3 în publicații de specialitate, documentație completă și actualizată, mijloace flexibile de configurare și de a înregistra date.

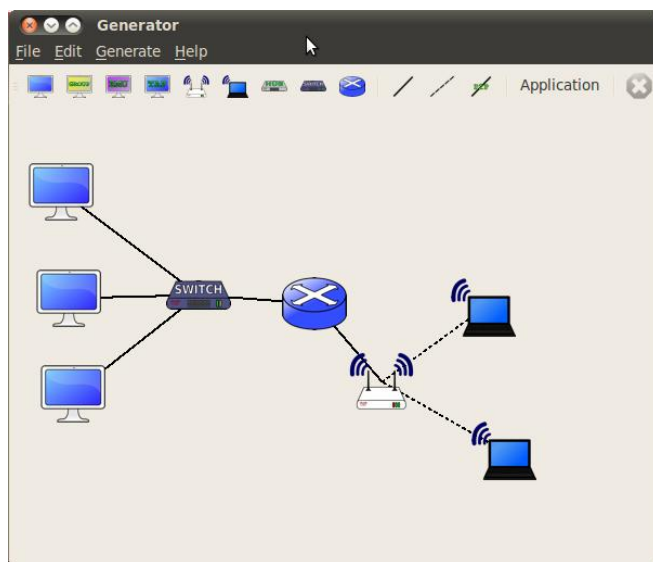


Figura 3.2: Simulatorul NS3

3.4.3 OPNET(Optimized Network Engineering Tool)

Este un software de simulare puternic și scump cu o mare varietate de posibilități de simulare de rețele heterogene cu protocoale diverse. A fost dezvoltat inițial la Massachusetts Institute of Technology (MIT) și din anul 1987 a devenit software comercial dezvoltat de către OPNET Technologies. Suporta modelarea rețelelor de comunicații cât și a sistemelor distribuite. Atât comportamentul cât și performanța sistemelor modelate pot fi analizate folosind simulări ale evenimentelor discrete.

Deoarece este un software comercial, OPNET oferă suport grafic și vizual foarte bun pentru utilizatori. Editorul grafic poate fi folosit pentru a construi topologii de rețea și entități începând cu nivelul aplicație până la nivelul fizic. Tehnici de programare orientate pe obiecte sunt folosite pentru a crea o mapare între designul grafic și implementarea sistemelor reale. De asemenea, parametrii pot fi ajustați și experimentele pot fi repetate cu ușurință prin intermediul interfeței grafice.

OPNET are trei funcții principale: modelarea, simularea și analiza. Pentru modelare, acesta oferă un mediu grafic intuitiv pentru a crea diverse modele de protocoale. Pentru simulare sunt folosite 3 tehnologii avansate și poate fi folosit pentru o mare varietate de studii. Pentru analiză, rezultatele simulării și datele pot fi analizate și afișate foarte ușor. Grafice, diagrame, statistici, chiar și animații pot fi generate cu ușurință de către OPNET pentru a fi convenabile utilizatorului. Începând cu 2008, OPNET Technologies a anunțat introducerea a doua aplicații capabile de managementul performanței. Aceste două noutăți includ capabilitatea de a vizualiza performanța aplicațiilor pentru organizații ce folosesc soluții de optimizare WAN precum și capabilitatea de a captura și vizualiza date NetFlow. Datorită efortului consistent, OPNET a devenit un simulator matur și este recunoscut în industrie. În plus, dezvoltarea OPNET a ținut cont de cerințele utilizatorilor și produsul este îmbunătățit constant, ceea ce îl face competitiv în comparație cu alte simulatoare de rețea comerciale.

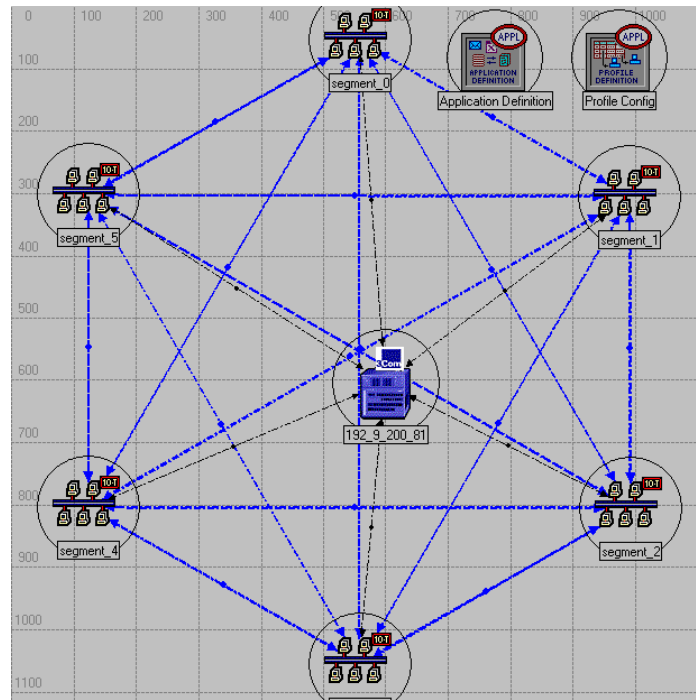


Figura 3.3: Simulatorul OPNET

3.4.4 OMNET++ (Optical Micro-Networks Plus Plus)

Este un framework pentru a dezvolta simulatoare de rețea, care conține o librărie extensibilă, modulară, bazată pe componente C++. Poate fi folosit pentru simularea rețelelor într-un sens mai larg precum rețele clasice, rețele wireless, rețele on-chip. Nodurile sunt programate în C++ și apoi asamblate într-un limbaj de nivel înalt numit NED (Network Description). Acest limbaj suportă specificarea de parametri variabili în descrierea rețelei, de exemplu numărul de noduri într-o rețea poate fi dinamic și poate fi configurat mai târziu în timpul execuției simulării. În acest caz, modulele care reprezintă nodurile sunt instanțiate dinamic de către simulator în timpul execuției. Această caracteristică este o consecință directă a designului orientat pe obiecte. OMNET++ formează o baza pentru a construi simulatoare orientate pe aplicații specifice.

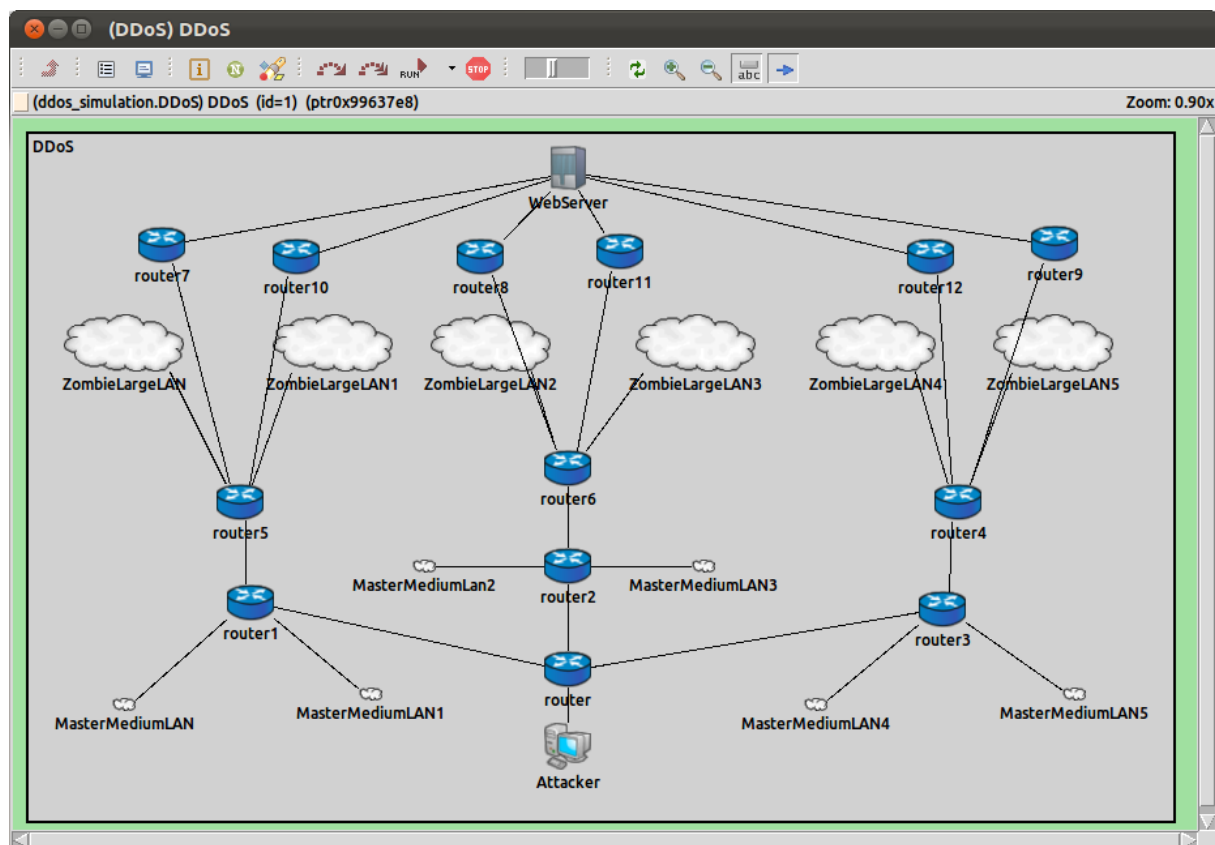


Figura 3.4: Simulatorul OMNET++

3.4.5 Cisco Packet Tracer

Packet Tracer este un software care poate fi utilizat în formare și educație, dar și în domeniul cercetării pentru simulări simple a rețelelor de calculatoare. Instrumentul este creat de Cisco Systems și este prevăzut pentru distribuție gratuită în facultate, pentru studenți și absolvenții care sunt sau au participat la programul Cisco Academy. Scopul Packet Tracer este de a oferi elevilor și profesorilor un instrument pentru a afla principiile de rețea precum și să dezvolte abilități specifice tehnologiei Cisco. Avantajul acestui simulator este interfața sa grafică ușor de folosit și faptul că este actualizat în mod constant, ultimele versiuni suportând Ipv6 (Internet Protocol version 6), BGP (Border Gateway Protocol) și HSRP(Hot Standby Router Protocol). Ca dezavantaj ar fi posibilitățile limitate de simulare, scopul acestuia fiind în principal didactic.

Acest software reproduce o mică parte din capacitățile hardware disponibile pe echipamentele Cisco cu propriul sistem de operare IOS. Totuși, există alte două aplicații software, GNS3 și Dynagen, care folosesc imagini reale de routere Cisco IOS. Diferența dintre cele două este că GNS3 are interfață grafică pe când Dynagen are doar linie de

comandă. Aceste două aplicații nu pot simula switch-uri și necesită multe resurse hardware, în special memorie RAM și putere de procesare.

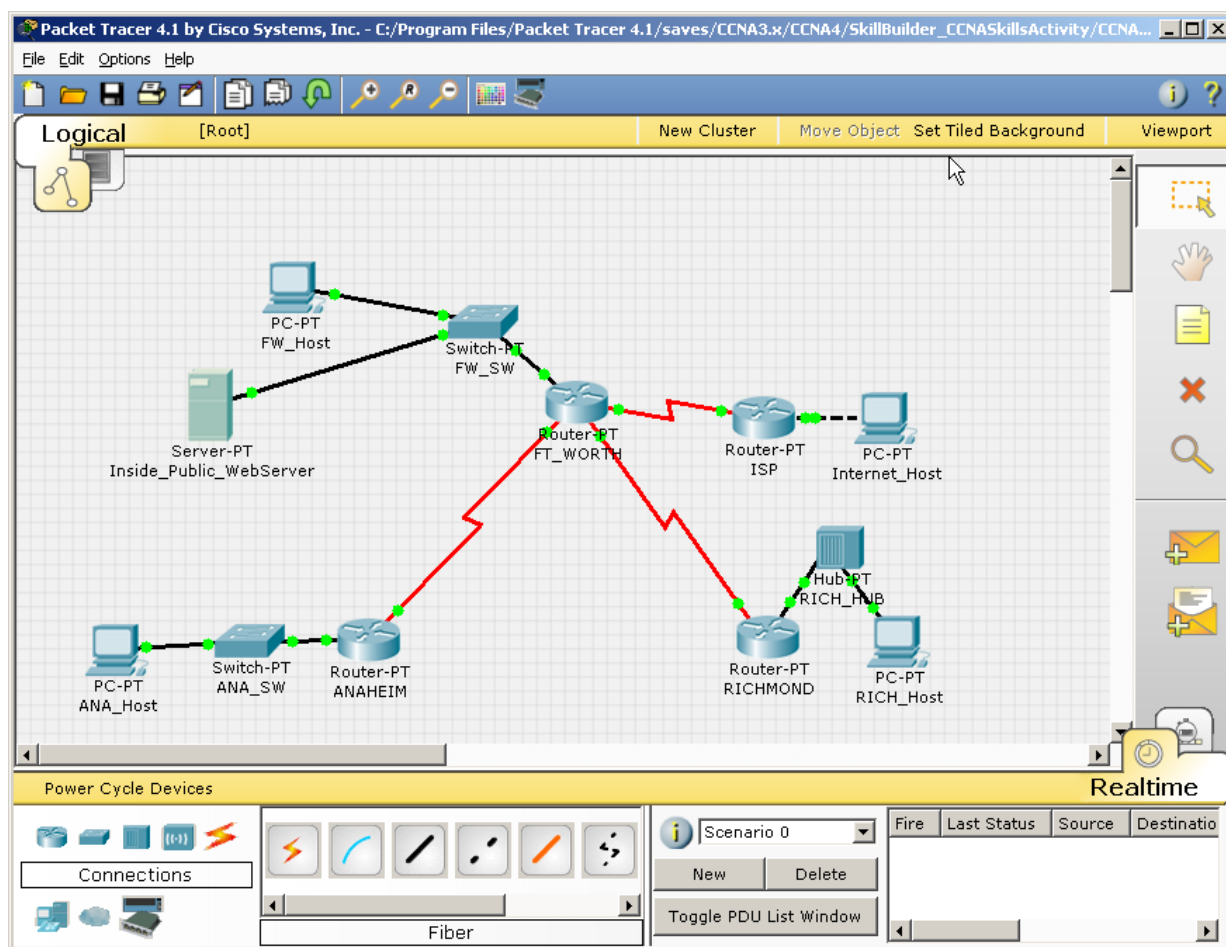


Figura 3.5: Cisco Packet Tracer

3.4.6 JiST(Java in Simulation Time)

JiST este un simulator de rețea care permite implementarea simulării în limbajul Java. Este, folosind în majoritatea cazurilor împreună cu SWANS, un simulator pentru rețele mobile și rețele ad hoc construit pe baza JiST. Simulările sunt compuse din entități care reprezintă elementele rețelei, de exemplu nodurile, iar evenimentele simulării sunt apelări ale metodelor asociate entităților. Entitățile sunt independente între ele, comunicând prin intermediul simulării de bază. În timp ce codul unei entități este un program Java obișnuit, doar interacțiunile dintre entități sunt relevante pentru timpul simulării. Așadar aceste interacțiuni între entități corespund momentelor de sincronizare și facilitează execuția în paralel a codului diferitelor entități de unde rezultă un potențial spor de performanță.

Pentru a executa implementarea în timpi specifici simulării, JiST folosește un încărcător de clase Java dinamic care rescrie codul aplicației în mod automat.

Din nefericire, dezvoltarea oficială JiST a fost oprită deoarece nu mai este întreținută de către autorul original, Rimon Barr. Cu toate acestea, o serie de îmbunătățiri și actualizări au fost lansate de către Universitatea Ulm.

Capitolul 4

4. Tehnologii WAN

4.1 ATM (Asynchronous Transfer Mode)

Utilizatorii LAN-ului sunt obișnuiți cu viteza și fiabilitatea propriilor rețele locale, și cum nevoile forțează LAN-urile să se extindă peste limitele lor originale, servicii precum liniile închiriate, ISDN, X.25, SMDS și Frame Relay au fost folosite pentru a interconecta LAN-uri. Dar utilizatorii se așteaptă la aceeași viteză și fiabilitate a traficului transmis peste rețele intermediare metropolitane și WAN. Aceste așteptări pun o substanțială presiune asupra furnizorilor de servicii și echipamente pentru a crea dispozitive care leagă diverse tehnologii de comunicație într-o rețea fără întreruperi. Această nevoie pentru viteză mare atât în LAN cât și în WAN este principala motivație pentru care au apărut primele aplicații ATM. Despre ATM se poate spune ca este o tehnologie de bază care încorporează alte tehnologii de LAN și WAN într-o singură arhitectură.

Numele tehnologiei ATM se justifică prin faptul că, în timp ce în rețelele telefonice majoritatea transmisiilor sunt sincrone (au legătură cu un semnal de ceas), în rețelele ATM transmisiile nu sunt sincrone. ATM a fost proiectat la începutul anilor 1990 și lansat la mijlocul acestei perioade incredibile [11].

ATM este un serviciu de comutare de celule (pachete de dimensiune fixă de 53 octeți). Un antet de 5 octeți indică destinația unei anumite celule și restul de 48 octeți sunt încărcătura efectivă. Câmpul de adresă este împărțit în două sub-câmpuri numite VCI (Virtual Channel Identifier) și VPI (Virtual Path Identifier), care formează o conexiune virtuală și au semnificație locală. Alte câmpuri importante din antet sunt CLP (Cell Loss Priority), care indică prioritatea unei celule în cazul în care traficul trebuie să fie aruncat. PTI (Payload Type Indicator) este un câmp de 3 biți dintre care primul bit indică faptul că celula conține date utile sau date pentru management și operații, al doilea bit este EFCI (Explicit Forward Congestion Indication), care indică dacă celula a întâlnit vreo congestie în drumul ei de la sursă la destinație și ultimul bit indică ultima celulă dintr-un bloc de celule. Câmpul HEC (Header Error Check) are rolul de a corecta un bit sau a detecta mai mulți biți eronați din primii 4 octeți ai antetului. Verificarea datelor utile este responsabilitatea unor protocoale de nivel superior.

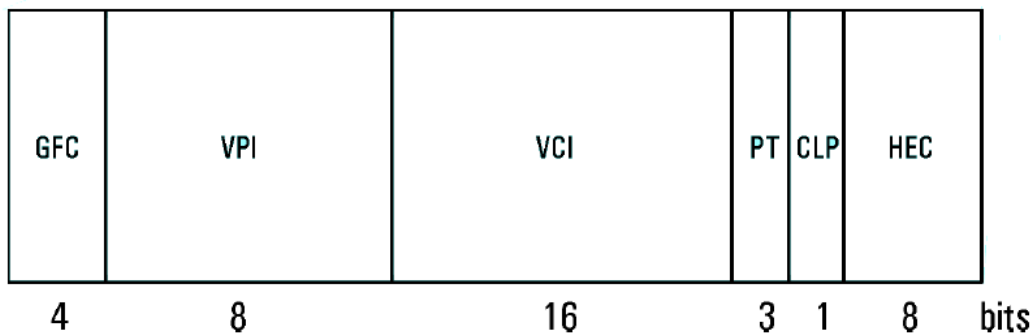


Figura 4.1: Antetul celulei ATM

ATM oferă posibilitatea creării unor conexiuni virtuale permanente cât și conexiuni virtuale temporare. Primele implementări ATM puteau crea doar conexiuni virtuale permanente dar ulterior au fost standardizate și conexiuni virtuale temporare. Traficul transportat prin rețele ATM trebuie divizat în celule mici cu încărcătura de 48 octeți. Acest proces se numește segmentare. Diverse metode de segmentare sunt utilizate pentru transmisii de voce, video sau date. Datorită vitezei ATM, celulele sunt comutate în hardware către linii cu lățimea de bandă mare. Celulele pot fi transmise cu orice viteză dar există unele standard: 45 Mb/s, 155 Mb/s, 622 Mb/s, 2.4Gb/s dar se poate atinge chiar și 10 Gb/s. Transmisia de celule de dimensiune mică și constantă în locul cadrelor de dimensiune variabilă a permis întrepătrunderea traficului sensibil la timp (ex: traficul telefonic) cu cadre de dimensiune mare (ex: traficul IP) cu un control mare al întârzierii transmisei. De aici rezultă și latențe mult mai mici față de alte tehnologii. Un alt avantaj este că celulele ATM pot satisface cerințe de lățime de bandă mult mai mari decât cele existente la ora actuală.

Deoarece rețelele ATM sunt orientate pe conexiune, transmisia datelor necesită mai întâi transmisia unui pachet pentru inițializarea conexiunii. Pe măsură ce pachetul de inițializare este transportat prin subrețea, toate ruterele de pe drumul pe care îl parcurge își creează câte o înregistrare în tabelele de dirijare în care înregistrează existența conexiunii și rezervă resursele necesare pentru ea. Conexiunile sunt de obicei denumite circuite virtuale, în analogie cu circuitele fizice utilizate în sistemele de telefonie. Majoritatea rețelilor ATM suportă și circuite virtuale permanente, care sunt conexiuni permanente între două gazde aflate la distanță.[*] Acestea sunt similare cu liniile închiriate din lumea telefoniei. Fiecare conexiune, indiferent dacă ea este temporară sau permanentă, are un identificator de conexiune unic.

Din momentul în care o conexiune a fost stabilită, se pot transmite date de oricare dintre cele două părți. O parte din antet reprezintă identificatorul de conexiune, astfel încât atât transmițătorul cât și receptorul, precum și toate ruterele intermediare pot ști corespondența dintre celule și conexiuni (care celule aparțin cărei conexiuni). Acest indentificator de conexiune permite fiecărui ruter să dirijeze fiecare celulă pe care o primește. Dirijarea celulelor este implementată direct în partea hardware a rutereleor și este

o operație rapidă. Justificarea alegerii de celule de dimensiune fixă este aceea că este mai ușor de construit partea hardware pentru dirijare dacă ea are de a face cu pachete scurte și egale ca dimensiune. Pachetele IP de lungime variabilă trebuie dirijate de programe (software), proces care este mai lent și acesta constituie cel mai mare avantaj al tehnologiei ATM. Un alt avantaj al rețelelor ATM este acela că partea hardware poate fi configurată cu ușurință să multiplice o celulă pe care o primește la intrare pe mai multe linii de ieșire, o proprietate obligatorie în cazul în care trebuie abordată transmisia unui program de televiziune difuzat către mai mulți receptori. Un lucru important este că celulele mici nu blochează nici o linie pentru prea mult timp, ceea ce face garantarea calității serviciilor (QOS – Quality Of Service) mai ușoară. Pentru că ATM realizează o conexiune, toate celulele urmează aceeași cale către destinație de unde rezultă faptul că este garantată ordinea lor, dacă se transmit celule într-o ordine, ele vor ajunge exact în aceeași ordine. Ceea ce nu se garantează însă este livrarea celulelor, nu se poate ști dacă odată trimise niște celule, ele ajung la destinație, acest lucru fiind de datoria protocoalelor de nivel superior, care repară eroarea cauzată de celulele pierdute.

4.1.1 Mecanismul de dirijare a traficului într-o rețea ATM

Există două tipuri de conexiuni ATM: căi virtuale (virtual paths - VP), care sunt identificate de către câmpul VPI și canale virtuale (virtual channels - VC), care sunt identificate de combinația dintre VPI și VCI. O cale virtuală este formată din mai multe canale virtuale, dintre care toate sunt comutate transparent în rețeaua ATM pe baza VPI-ului comun. O cale de transmisie este formată din mai multe VP-uri. Asemănător cu o linie telefonică, cu ajutorul căilor virtuale se pot direcționa mai eficient canalele virtuale. De-a lungul unei căi, dacă un switch cedează, circuitele pot fi re-direcționate modificând calea virtuală, în loc de a re-direcționa fiecare canal individual. Figura 1.3 ilustrează modul în care canalele virtuale se concatenează pentru a crea căi virtuale, care, la rândul lor se concatenează pentru a crea o cale de transmisie.

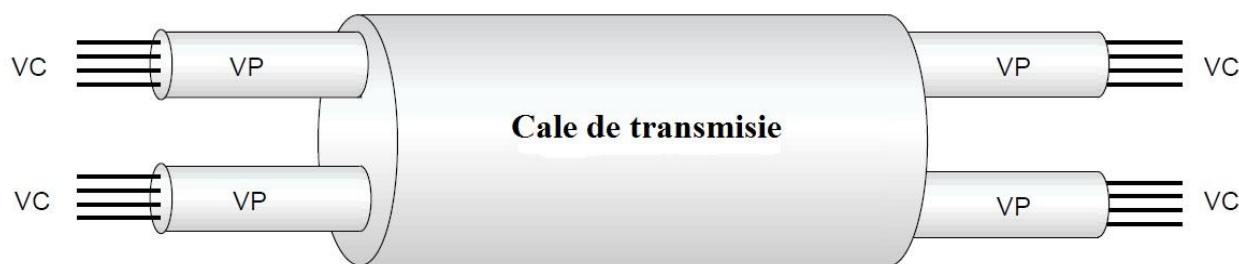


Figura 4.2: Căile virtuale sunt formate din canale virtuale

Funcționarea unui switch ATM este destul de simplă: o celulă este recepționată pe o interfață cu o valoare VCI sau VPI cunoscută. Switch-ul caută valoarea conexiunii într-un tabel local de translații pentru a determina portul (sau porturile) de ieșire ale conexiunii și noua valoare VPI/VCI a conexiunii de pe portul respectiv. După aceasta, switch-ul retransmite celula pe portul (sau porturile) de ieșire cu identificatorii de conexiune potriviți. Deoarece toate valorile VCI și VPI au semnificație locală pe o anumită conexiune, aceste valori sunt modificate la fiecare switch, după cum este nevoie.

4.1.2 Segmentarea și reasamblarea

Pachetele folosite de către stiva de protocoale TCP/IP sunt aproape întotdeauna mai mari decât celulele ATM. Pentru ca dispozitivele ATM să comunice cu aceste rețele existente, pachetele trebuie segmentate pentru a fi de dimensiunea celulelor înainte de a intra în rețea ATM și apoi reasamblate în pachete mai mari când părăsesc rețeaua ATM. În continuare sunt prezentați pașii necesari pentru segmentare și reasamblare:

- Un nod ATM primește un cadru de date dintr-un nivel superior.
- Se înlătură CRC-ul (Cyclic Redundancy Check) pachetului. Aceasta este porțiunea din pachet care se ocupă cu detectarea erorilor.
- Se adaugă un antet și subsol pachetului care conține informații despre nodul terminal ATM.
- Se adaugă date pentru ca dimensiunea pachetului să fie un multiplu de 32 biți.
- Se adaugă un nou CRC pachetului.
- Se segmentează pachetul în segmente de 48 biți (care includ antetul și subsolul)
- Se plasează fiecare segment de 48 biți într-o celulă ATM și se adaugă un antet care include VCI-ul.
- Switch-urile ATM de-a lungul canalului se uită doar la VCI și la portul pe care a sosit celula, se consultă cu un tabel de direcționare care conține porturile de intrare și porturile de ieșire și trimit celula mai departe.
- La nodul ATM terminal, celulele sunt reasamblate în pachete folosind pașii anteriori în ordine inversă.
- După ce un pachet este reasamblat, este verificat de erori de un protocol de nivel superior.

4.2 MPLS (Multiprotocol Label Switching)

Comutarea Multiprotocol cu Etichete (Multi Protocol Label Switching - MPLS) reprezintă o arhitectură în care nodurile terminale adaugă o etichetă unui pachet IP ce identifică drumul spre destinație iar pachetele sunt direcționate pe baza etichetei fără inspectarea antetului inițial.

MPLS folosește o soluție ce integrează atât controlul rutării IP (nivelul 3 din modelul OSI), cât și comutarea de la nivelul legăturii de date (nivelul 2 din modelul OSI), de aceea se spune că MPLS este un „protocol de nivel 2,5”. Această tehnologie oferă bazele unor servicii de rutare avansate, rezolvând o serie de probleme:

- Se adresează problemelor privind scalabilitatea, legate de modelul IP-over-ATM;
- Reduce complexitatea operațiilor din rețea;
- Facilitează apariția de noi posibilități de rutare, ce îmbunătățesc tehnicile de rutare IP existente;
- Oferă o soluție standardizată, ce are avantajul interoperabilității între diverși furnizori de produse și servicii.

Într-o rețea tradițională IP fiecare router efectuează o căutare a adresei de destinație, determină next-hop-ul pe baza tabelului de rutare și trimite pachetul către acel next-hop. Aceste operații au loc în fiecare router, fiecare având propriile sale decizii independente de rutare, până când este atinsă destinația. Într-o rețea MPLS doar primul echipament efectuează o căutare a adresei destinație dar, în loc să găsească un next-hop, găsește router-ul destinație și o cale predefinită către acel router destinație. Următorul pas este aplicarea unui etichete pe baza informațiilor găsite și următoarele routere folosesc această etichetă pentru a direcționa traficul fără a mai fi nevoie de a efectua alte căutări specifice IP. Când pachetul ajunge la router-ul destinație această etichetă este înlăturată și pachetul este direcționat în continuare folosind IP tradițional.

4.2.1 Concepte de bază MPLS

Cale comutată după eticheta (Label Switched Path - LSP).

Este unul dintre cele mai importante concepte folosite în practică. În esență, LSP-ul, este un tunel unidirecțional între o pereche de routere, care este dirijat în interiorul unei rețele MPLS. Un LSP este necesar pentru a dirija orice trafic. Fără un LSP configurat nu putem avea trafic.

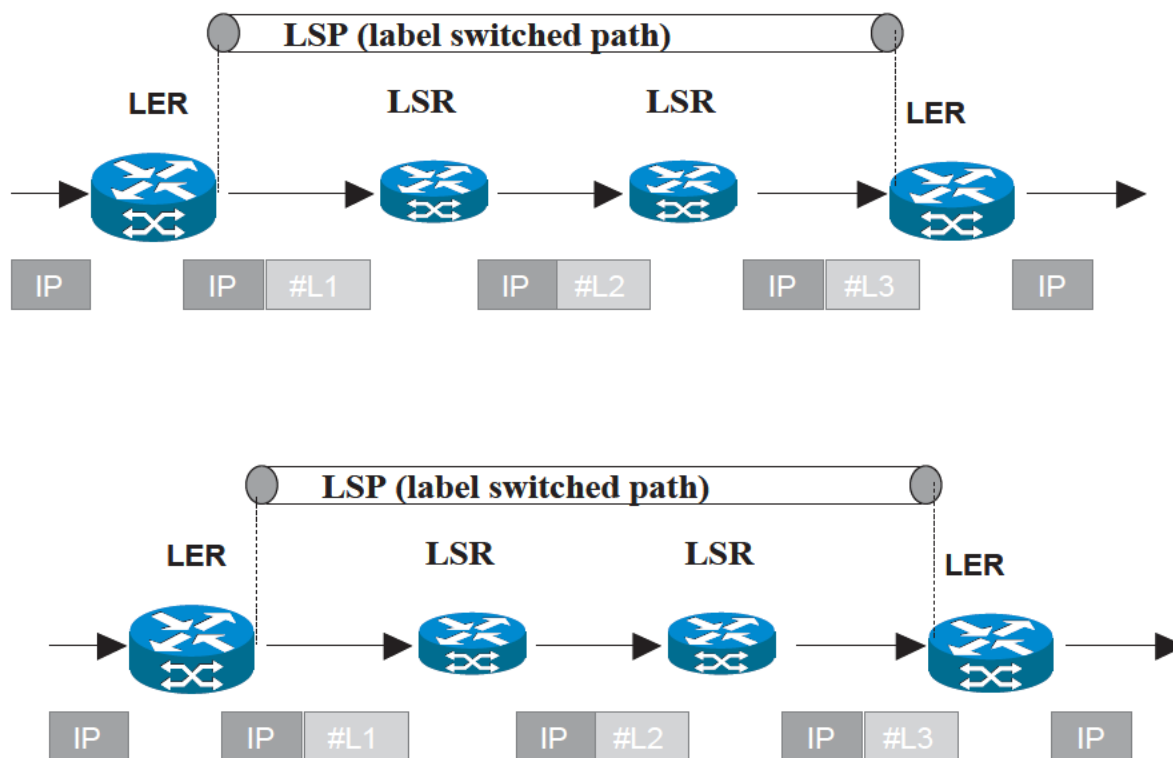


Figura 4.3: Cale comutată după etichetă (Label Switched Path - LSP)

Router de etichetare marginal (Label Edge Router - LER). Se mai numește nod de intrare (ingress node), deoarece este router-ul care încapsulează prima dată un pachet într-un LSP MPLS. De asemenea, este router-ul care face alegerea inițială a LSP-ului.

Router de comutare a etichetelor (Label Switching Router - LSR). Este un router care efectuează doar comutare MPLS pe baza etichetei de-a lungul unui LSP.

Nod de ieșire (egress node). Este router-ul situat la capătul unui LSP și este cel care înlătură eticheta.

Legătura protejată (Protected Link). Este o legătură fizică cu o formă de redundanță integrată astfel încât transferurile de date nu sunt întrerupte de căderea unei componente a legăturii. O legătură protejată apare în nivelul de control MPLS ca un singur punct de conexiune în interiorul rețelei. Există multe scheme de legături protejate, dar una populară folosește protocoale SONET/SDH pe o buclă de fibră optică.

Legături paralele (Parallel or Bundled Links). Există mai mult de o legătură între o pereche de noduri într-o rețea. Spre deosebire de o legătură protejată, aceste legături individuale apar ca puncte de conexiune separate în rețea. Pot fi organizate ca

entități distincte suportând rute diferite (dar paralele) în interiorul rețelei sau pot fi organizate ca o singură legătură unde alegerea componentei individuale este făcută doar de către nodurile conectate direct la legătură.

Cale alternativă (Alternate Path). Este o rută diferită în interiorul rețelei care leagă aceleași puncte terminale. Legăturile paralele asigură cele mai simple căi alternative. Căi alternative mai complicate trebuie să traverseze legături distincte și noduri diferite.

Clasa de dirijare echivalentă (Forward Equivalent Class - FEC). Reprezintă un set de pachete care sunt tratate identic de către LSR. Astfel, un FEC este un grup de pachete IP care sunt dirijate peste același LSP, sunt tratate la fel și pot fi asociate unei singure etichete de către un LSR chiar dacă pachetele sunt diferite în privința informației din antetul de nivel rețea. Eticheta minimizează informațiile esențiale despre pachet. Acestea pot include destinația, precedentă, informații legate de QoS chiar și întreaga rută a pachetului, după cum decide LSR-ul de intrare bazat pe politici administrative.

Eticheta (Label sau shim label). Reprezintă un identificator scurt, de lungime fixă și cu valoare locală al unui FEC. Deși informațiile de nivel rețea sunt folosite pentru asocierea etichetei, eticheta nu conține în mod direct informații în nivel rețea precum adresa sursă și destinația. Etichetele au doar valoare locală, ce înseamnă că o etichetă este utilă și relevantă pe o singură legătură între LSR-uri adiacente.

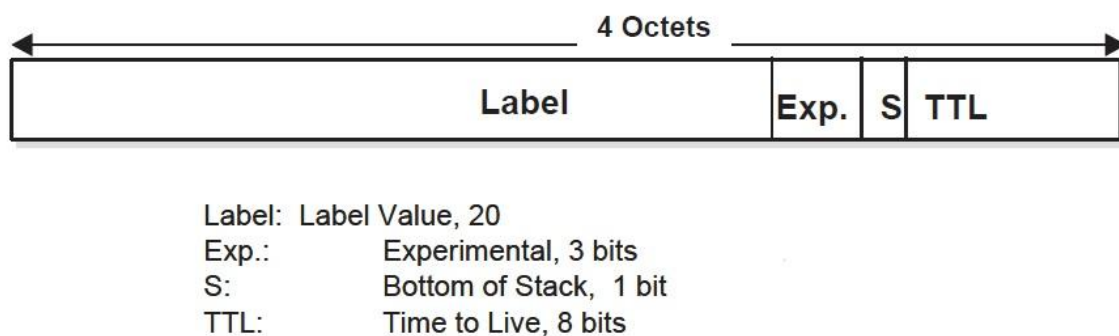


Figura 4.4: Structura antetului MPLS

Pentru a folosi un LSP, acesta trebuie să fie anunțat către celelalte routere. Un LSP este un tunel întins pe o întreaga rețea, dar o etichetă are valoare doar locală de-a lungul unei legături. Un protocol de semnalizare MPLS asociază LSP-uri cu anumite valori ale etichetelor. Sunt folosite două protocoale principale pentru semnalizare MPLS:

- Protocol de distribuție a etichetelor (Label Distribution Protocol - LDP). Este un protocol simplu, fără constrângeri (nu suportă ingineria traficului)
- Protocol de rezervare a resurselor cu ingineria traficului (Resource Reservation Protocol with Traffic Engineering – RSVP-TE). Este un protocol mai complex,

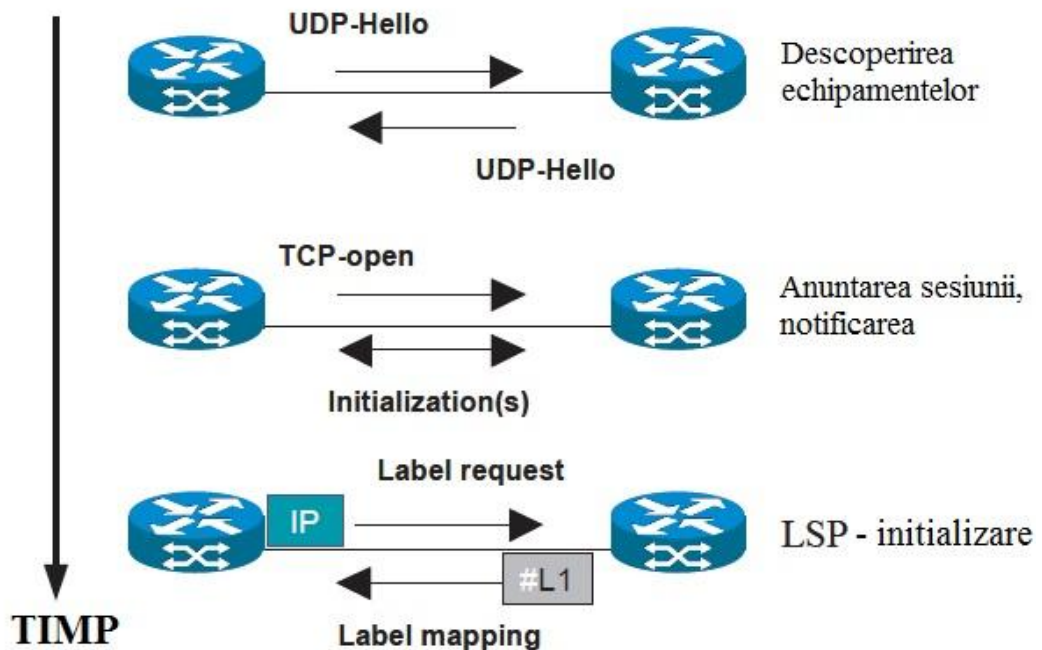


Figura 4.5: Funcționarea LDP (Label Discovery Protocol)

Anumite rețele complexe folosesc ambele protocoale în același timp. LDP se folosește de obicei pentru servicii VPN (Virtual Private Network) pentru transport de date. Multe rețele MPLS configurează tunele ce folosesc LDP în interiorul protocolului RSVP. Etichetele pot fi suprapuse una peste alta și eticheta de deasupra este folosită pentru a controla livrarea pachetului. Când este atinsă destinația, eticheta de deasupra este înlăturată și următoarea etichetă preia sarcina de a controla livrarea pachetului mai departe.

Există două feluri de a termina un LSP, Null implicit și Null explicit. Null implicit este o tehnică de optimizare și înseamnă că eticheta este înlăturată de către penultimul router, nu de către ultimul router ca în mod obișnuit. Acest lucru reprezintă o optimizare deoarece ultimul router nu mai este nevoie să caute calea de dirijare și să înlătore eticheta,

fiind necesară doar o dirijare pe baza informației continute. Null explicit înseamnă că eticheta este păstrată până la ultimul router.

IGP-urile (Interior Gateway Protocol) clasice folosesc rutare fără ingineria traficului, de exemplu o metrică (cost) pe o legătură și un algoritm de cea mai scurtă cale (Shortest Path First - SPF) pentru a găsi cea mai scurtă cale. Ingineria traficului adaugă la acestea anumite constrângeri, de exemplu găsirea drumului cel mai scurt și care are lățime de bandă disponibilă. Acest lucru se mai numește rutare cu constrângeri, folosind un algoritm de cea mai scurtă cale cu constrângeri (Constraint Shortest Path First - CSPF). Principiul ingineriei traficului este simplu: este mai bine să alegem o cale necongestionată chiar dacă întârzierea poate fi mai mare, decât să congestionăm cea mai scurtă cale în timp ce pe alta legătura avem lățime de bandă disponibilă și rămâne nefolosită.

Ingineria traficului se poate realiza și manual, modificând costurile IGP-ului, dar acest lucru este scalabil doar până la un anumit punct. Pe măsură ce rețelele devin din ce în ce mai complexe, acest lucru devine mai greu de întreținut și modificarea unui cost IGP chiar și cu o unitate poate afecta dirijarea traficului la o distanță de zeci de hop-uri. De aceea nu este recomandată ingineria traficului modificând manual costurile IGP-ului.

4.2.2 Potențiale eșecuri ale resurselor

Orice resursă a rețelei poate eșua. Pentru a asigura o rețea cu high availability, furnizorul rețelei trebuie să prezică aceste căderi și să planifice din avans.

Eșecul tradițional este căderea unei legături cauzată de tăierea unui cablu sau scoaterea din soclu. Deasemenea, eșecul este posibil și în rețele optice rezultând o lipsă a luminii, ca urmare a defectării laserului. Alte probleme obișnuite includ defectarea routerelor când un întreg LSR eșuează. Acest lucru poate fi cauzat de o întrerupere de curent sau o defectare a unei componente esențiale a LSR-ului. Când un router eșuează, toate legăturile direct conectate eșuează.

Dacă are loc un eșec hardware al unei componente, LSR-ul își poate reveni. Eșecul unei plăci de rețea (sau port) este echivalent cu eșecul unei legături, dar poate fi recuperată intern dacă există o placă de rețea de rezervă pentru acea legătură. În acest mod, o placă de rețea de rezervă poate fi privită ca un element al unei legături protejate. Ambele plăci de rețea sunt conectate cu legătura și doar cea principală trimite date, iar când aceasta eșuează, placa de rezervă este deja conectată cu legătura.

4.2.3 Obiective pentru supraviețuirea în caz de eșec

Obiectivul de bază pentru supraviețuirea în caz de eșec este minimizarea întreruperii traficului de date a oricărui tip de eșec.

Pe cât posibil, LSP-urile prestabilite să nu fie întrerupte deloc în timp ce rețeaua se reface după eșec. Acest lucru înseamnă că legăturile ar trebui menținute și pachetele nu trebuie aruncate. În practică, multe eșecuri întrerup traficul existent în timp ce se face o tranziție de la vechile resurse la cele noi, dar aceste întreruperi ar trebui să fie minimizate pe cât posibil. O cifră de 60ms este folosită în lumea telecomunicațiilor ca fiind cea mai mare întârziere a traficului de voce care poate fi acceptată de creierul uman până când devine ca o pauză care oprește înțelegerea fluxului de voce. Acest lucru înseamnă că pentru traficul de voce, în mod ideal, orice eșec trebuie detectat, raportat și reparat în timp de 60ms.

Chiar dacă repararea LSP-ului durează mai mult de 60ms este important să se restabilească conexiunea în mod automat. Să considerăm un utilizator telefonic, ideal ar fi să nu observe eșecul deloc, dar este mai bine să audă câteva clic-uri decât să piardă de tot conexiunea și să fie nevoit să sune din nou. Dacă este întrerupt traficul de date, trebuie făcute anumite considerații dacă s-au pierdut date și câte s-au pierdut. Nici rețelele IP, nici alte rețele precum ATM sau Frame Relay nu încearcă să asigure livrarea datelor, decât prin folosirea altor protocoale de nivel superior precum TCP. Totuși, dacă s-au pierdut prea multe date, astfel de protocoale pot declara conexiunea eșuată și să încerce reconectarea.

Un alt obiectiv cu prioritate ceva mai mică este acela că, serviciul de semnalizare ar trebui să rămână funcțional. Altfel spus, ar trebui să fie posibilă stabilirea de noi conexiuni după ce eșecul a fost reparat. Deși nu este de dorit, în general, este acceptat ca un utilizator să reîncece inițializarea unei conexiuni dacă conexiunea nu s-a stabilit de prima dată. Având în vedere date statistice cu privire la posibilitatea inițierii unei conexiuni în timp ce un eșec este reparat, de multe ori este considerat acceptabil ca semnalizarea să fie suspendată temporar.

Procesul reparării într-o anumită parte a rețelei ar trebui să întrerupă cât mai puțin alte părți ale rețelei. Distribuind informații despre eșec în toată rețeaua ar putea întrerupe alte semnalizări și trafic de date. O soluție de reparare nu are nici o valoare dacă este proiectată să supraviețuiască unui singur eșec. Mulți furnizori de rețele și echipamente nu încearcă să asigure soluții care să supraviețuiască mai multor eșecuri simultane. Chiar dacă acest lucru reduce complexitatea, implică faptul că timpul de recuperare al unui

echipament trebuie să fie scurt pentru a asigura că perioada de vulnerabilitate este scurtată pe cât posibil.

Toate soluțiile acestor cerințe implică forme de redundanță precum legături paralele, legături suplimentare în rețea sau prin adăugarea de componente hardware în interiorul switch-ului. Costul acestor soluții impune o cerință adițională conform căreia resursele redundante ar trebui să fie minime și preferabil să fie partajate de utilizatori multipli.

4.2.4 Avantaje MPLS

Deoarece MPLS folosește dirijarea pe baza etichetei de dimensiune fixă, comutarea fiecărui pachet este determinată de o singură cautare indexată într-un tabel de comutare. Acest lucru simplifică funcționarea routerului în comparație cu algoritmul de găsimă al celui mai lung prefix folosit pentru dirijarea datagramelor normale.

Unul din principalele avantaje MPLS este abilitatea sa de a realiza ingineria traficului în rețele IP. Ingineria traficului este necesară pentru a se asigura că traficul este dirijat printr-o anumită rețea în cel mai eficient și sigur mod posibil. Cu ajutorul ingineriei traficului se poate distribui traficul în toată infrastructura rețelei și permite furnizorilor de internet să dirijeze traficul într-un mod în care pot oferi cele mai bune servicii utilizatorilor în termeni de rate de transfer și întârzieri.

Dirijarea QoS bazată pe sursă este un mecanism de dirijare în care LSR-urile pe care le va străbate o cale sunt determinate în nodul sursă pe baza resurselor disponibile în rețea și pe baza cerințelor QoS. Cu alte cuvinte, este un protocol de rutare care și-a extins criteriile de selecție a căilor să includă parametrii QoS precum lățimea de bandă disponibilă, gradul de utilizare a căii și a legăturii, resursele consumate în nod, întârziere și latentă.

O rețea privată virtuală (Virtual Private Network - VPN) bazată pe internet folosește infrastructura deschisă și distribuită a internetului pentru a transmite date între locații, menținând securitatea prin folosirea unui protocol de încapsulare pentru a stabili tuneluri. O rețea privată virtuală poate fi asemănată cu un sistem de linii deținute sau închiriate care pot fi folosite de o singură companie. Scopul principal al unui VPN este să ofere companiilor aceleași capabilități ca liniile închiriate private cu un cost mult mai mic prin folosirea infrastructurii publice. Arhitectura MPLS satisface toate cerințele necesare pentru suportul VPN prin stabilirea de tuneluri LSP folosind rutare explicită. Așadar, tehnologia MPLS folosind tuneluri LSP permite furnizorilor de servicii să livreze aceleași capabilități utilizând aceeași infrastructură cu care livrează și servicii de internet.

Comutarea ierarhică se face pe baza suprapunerii etichetelor (sau controlul pachetelor pe mai multe nivele) și astfel se permite încapsularea unui LSP în interiorul

altui LSP. Acest lucru nu este o noutate în tehnologia rețelelor, ATM furnizează comutare ierarhică pe două nivele cu noțiunea de cale virtuală și circuit virtual. Totuși, MPLS permite suprapunerea arbitrară a LSP-urilor, furnizând astfel mai multe nivele de control a pachetelor.

Comutarea etichetelor permite o separație mai completă între rutarea în interiorul domeniului și rutarea între domenii, fapt care ajută la îmbunătățirii scalabilității procesului de rutare. Mai mult, scalabilitatea MPLS este îmbunătățită de FEC deoarece se agregă fluxurile și suprapunerea etichetelor pentru combinarea LSP-urilor. Asocierea unei etichete pentru fluxuri de date individuale nu este ideea dorită pentru scalabilitate deoarece crește numărul de etichete, care în consecință determină creșterea tabelului LIB. Pentru că FEC permite agregarea fluxurilor se îmbunătățește scalabilitatea MPLS. În plus, LSP-uri multiple, ce aparțin de FEC-uri diferite pot fi unite într-un singur LSP.

Capitolul 5

5. Soluții de recuperare MPLS

În ultimii ani, noi servicii și aplicații au fost dezvoltate cu caracteristici de timp real și orientate pe conexiune, servicii precum Voice-over-IP sau protocolul RTSP (Real Time Streaming Protocol).

Căderea unei legături majore sau a unui router important poate avea efecte severe asupra acestor servicii și protocoale. După terminarea redirijării, aceste servicii ar putea suferi de o degradare a serviciului calității, deoarece ruta alternativă ar putea fi mai lungă sau mai aglomerată. De luat în considerare că traficul care nu este afectat în mod direct de eșec dar redirecționat pe o cale alternativă este și el afectat de această degradare.

Pe de altă parte, durata întreruperii datorită căderii unei legături sau căderii unui nod este, de obicei, prea mare pentru servicii în timp real și aplicații multimedia ca să își păstreze sesiunile. În același timp, fluxurile QoS pot suferi de o reducere a calității pe rutele alternative și prin urmare nu mai pot fi restabilite.

Deoarece MPLS este o arhitectură orientată pe conexiune, în caz de eșec, trebuie să se stabilească un nou LSP și după aceea să comute pachetele de la punctul de eșec sau de la alt nod către LSP-ul nou format. Din această cauză MPLS are un răspuns de restaurare încet ca urmare a unui eșec al unei legături sau al unui nod dintr-un LSP.

Una din provocările unui protocol orientat, precum MPLS este garantarea serviciului în timpul eșecului. Abilitatea de a redirija traficul rapid în jurul unui punct de eșec sau în jurul unui punct de congestie într-un LSP poate fi important pentru rețele MPLS astfel încât să asigure calitatea serviciului chiar și în situații de eșec.

Mai multe metode au fost propuse pentru a redirija traficul în MPLS. Există două scheme luate în considerare în cadrul IETF oferind abordări diferite a problemelor de restaurare a LSP-ului în rețele bazate pe MPLS. Primul este cel mai rapid mecanism de redirijare disponibil care a fost propus de către Haskin și Krishnan [9], iar cel de-al doilea este un mecanism mai lent dar mai puțin complex propus de Makam et al. [8] cunoscut ca tunel de back-up bazat pe RSVP (RSVP-based Backup Tunnel).

5.1 Redirijarea

Redirijarea este o tehnică care poate fi folosită și în rețele cu comutare de etichete și rețele cu comutare de pachete. Redirijarea este definită ca crearea unui nou traseu sau a unui segment al traseului la cerere pentru restaurarea traficului după apariția unui eșec. Așadar este un mecanism de recuperare în care traseul de revenire sau segmentul traseului este creat dinamic după detecția eșecului pe traseul în funcțiune. Pentru acest scop, un traseu alternativ sau un traseu de rezervă sunt necesare în afară de calea principală folosită de trafic. Traseul principal și traseul de rezervă ar trebui să fie total disjuncte. Componentele unui rețea sunt formate, în mare parte, din legături și noduri. Deoarece eșecul unui nod cauzează căderea legăturilor adiacente nodului, vom folosi căderea legăturilor ca un eșec al rețelei.

Atunci când are loc o cădere a unei legături pe traseul principal, procesul de restaurare începe în mod automat. Pașii principali care trebuie efectuați de către metoda de redirijare sunt detectarea eșecului, notificarea eșecului, calcularea traseului alternativ și redirijarea traficului de pe traseul principal pe traseul alternativ.

5.2 Modelul lui Makam

Una din primele modele propuse de recuperare MPLS a fost prezentată de Makam în draftul [OSMH01]. Modelul propus este referențiat ca modelul lui Makam și furnizează protecție capăt la capăt pentru un LSP prin instalarea unui traseu de recuperare globală între LSR-ul de intrare și LSR-ul de ieșire. Acest traseu de recuperare este total disjunct față de traseul principal. Atunci când este detectat un eșec oriunde de-a lungul căii principale, un semnal de indicare a eșecului (fault indication signal - FIS) este folosit pentru a transporta informații despre eșec către nodul responsabil cu comutarea traseelor (Path Switch LSR - PSL). Când acest semnal este primit, PSL-ul comută traficul către traseul de recuperare.

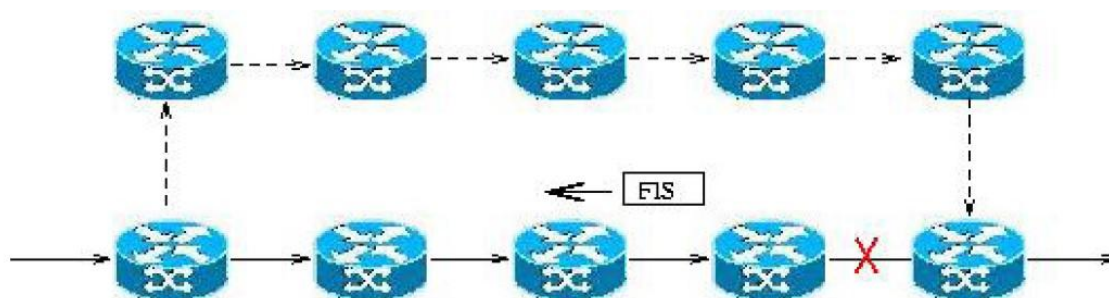


Figura 5.1: Modelul lui Makam

Figura 5.1 arată cum semnalul FIS este trimis în direcția opusă traseului principal atunci când un eșec este detectat pe acest traseu.

Acest model are propuneri atât pentru un traseu de recuperare prestabilit cât și pentru un traseu de recuperare stabilit dinamic (redirijat). Deoarece traseul de recuperare stabilit dinamic va adauga mai mult timp operației de restaurare, traseul de recuperare prestabilit este folosit de cele mai multe ori cu modelul lui Makam[8].

5.3 Modelul lui Haskin (Reverse backup)

Când este folosită recuperarea globală ca în modelul lui Makam, PSL-ul trebuie informat despre eșecul în traseul principal înainte ca traficul să fie comutat pe calea de recuperare. Când acest lucru este realizat cu ajutorul semnalului FIS, traficul va fi trimis în continuare pe traseul principal până când semnalul va fi recepționat de către PSL. Înseamnă că vor fi pierdute pachete la LSR-ul care a detectat eșecul deoarece acest nod nu are nici o informație despre cum să dirijeze pachetele. Dacă eșecul este situat la distanță mare de punctul de reparare și rata de transmisie este mare, numărul de pachete pierdute poate fi foarte mare.

Idea din spatele recuperării inversate este redirijarea traficului în direcția pe care a sosit la punctul de eșec de pe traseul principal către PSL. Din moment ce un LSR detectează un eșec pe traseul principal, acesta redirijează traficul pe LSP-ul alternativ care a fost creat în sens opus traseului principal. Atunci când traficul inversat este recepționat de PSL, acesta îl comută pe o cale de protecție globală. Atât calea inversată cât și traseul de protecție globală sunt prestabilite. Această schemă a fost introdusă de către Haskin în draftul [HK00] și este referențiat ca modelul lui Haskin.

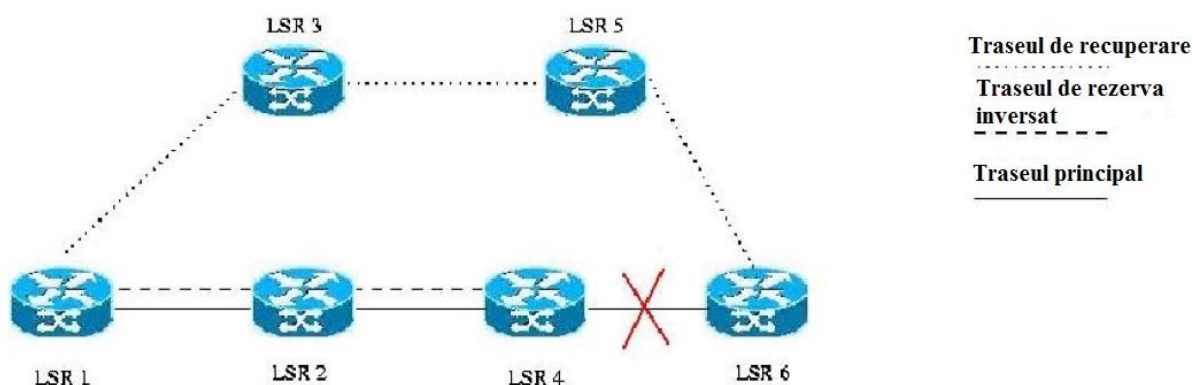


Figura 5.2: Modelul lui Haskin

Figura 5.2 arată un exemplu de recuperare inversată. Traseul principal și traseul de recuperare sunt stabilite la fel ca și în protecția globală, dar în plus există o cale inversată de la LSR4 la nodul de intrare LSR1. Atunci când un eșec este detectat de LSR4, traficul este redirijat înapoi către LSR1 prin LSP-ul de rezervă inversat și apoi comutat pe calea globală de recuperare. După cum se observă, acest model protejează împotriva eșecurilor situate oriunde pe calea de la LSR1 la LSR6.

Există și opțiunea de a face PSL-ul să comute traficul direct pe traseul global de recuperare, din moment ce se va observa că este folosit traseul de rezervă inversat. În acest fel, traficul inversat este folosit ca un semnal FIS, când primul pachet este recepționat de pe traseul de rezervă inversat, PSL-ul începe să comute traficul direct pe calea de recuperare globală. Astfel, traseul de recuperare se scurtează, deoarece nu mai este nevoie să comutăm traficul pe traseul principal care este eșuat.

Această optimizare are un dezavantaj, până când PSL-ul recepționează traficul inversat, pachetele vor fi trimise în continuare pe traseul principal care a eșuat. Atunci când PSL-ul primește trafic de pe traseul de rezervă inversat, acesta începe să comute traficul pe calea globală de restaurare. Pentru o perioadă scurtă de timp, traficul intrat va fi combinat cu traficul inversat pe traseul de recuperare. Această reordonare de pachete poate cauza probleme protocoalelor de nivel superior care au nevoie ca pachetele să ajungă în ordine.

O altă problemă cu modelul lui Haskin este utilizarea mai puțin eficientă a resurselor, deoarece traseul de recuperare este mai lung decât traseul principal. Partea pozitivă este numărul scăzut de pachete pierdute când are loc un eșec[9].

5.4 Redirijarea rapidă (Fast Rerouting)

Mecanismul de restaurare Fast Rerouting prestabilește calea alternativă la calea protejată înainte să aibă loc orice eșec. Criteriul creării căii alternative prestabilite este bazat pe politici de rutare, cerințe de restaurare a traficului protejat și considerații administrative. Când are loc un eșec, LSR-ul responsabil cu redirijarea traficului comută traficul de pe calea protejată pe calea alternativă prestabilită. Deoarece modelul restaurării prestabilește calea de recuperare înainte de apariția unui eșec, timpul de recuperare este mai mic decât modelul redirijării.

Deoarece întregul traseu este setat într-un LSP, nu pot apărea bucle de rutare în timpul convergenței, chiar dacă traseul este suboptimal.

Avantajul major al modelului de redirijare rapidă este timpul redus de recuperare, deoarece nu este nevoie de calculul unui nou traseu în caz de eșec. Rerutarea rapidă este folosită ca un mecanism de recuperare locală care folosește protecția legăturii sau a nodului prin stabilirea unui traseu alternativ care utilizează comutarea pentru a ocoli legătura căzută sau nodul eșuat. Dacă se folosește redirijarea rapidă de la un capăt la altul, atunci traseele de recuperare trebuie să fie prestabilite pentru fiecare legătura sau fiecare nod de pe traseul principal. În [PSA00] sunt definite extensii ale protocolului RSVP-TE pentru stabilirea unui LSP cu tunele de rezervă de la un capăt la altul. Sunt descrise două tehnici în draft, modelul de rezervă unu la unu(one-to-one backup) și modelul de rezervă cu facilitare(facility backup)[10].

5.4.1 Redirijarea rapidă cu rezervă unu la unu

Folosind tehnica de redirijare rapidă cu rezervă unu la unu, se calculează un LSP de rezervă pentru fiecare LSR din calea protejată. Aceste LSP-uri de rezervă sunt configurate să folosească restaurarea nodurilor și dacă nu este posibil, să folosească restaurarea legăturilor. Pentru a proteja în întregime un LSP care traversează N noduri pot fi până la $N-1$ LSP-uri de rezervă.

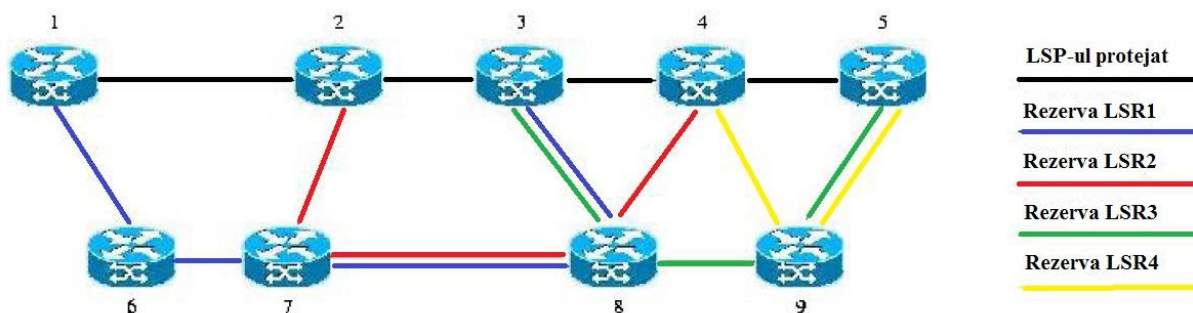


Figura 5.3: Redirijarea rapidă cu rezervă unu la unu

În figura de mai sus, LSP-ul protejat are 4 LSP de rezervă. Alternativele 1-3 folosesc protecția nodului, în timp ce alternativa 4 folosește protecția legăturii. Dacă are loc un eșec oriunde pe calea protejată, LSR-ul care a detectat căderea poate oricând să comute traficul pe o cale alternativă fără a mai fi nevoie să trimită un semnal FIS pentru ca restaurarea să înceapă. Prin urmare, operațiunea de recuperare devine o decizie locală pentru LSR-ul care a detectat eșecul. Folosind acest model, există câteva extensii sugerate pentru protocolul RSVP-TE care fac posibilă găsirea și crearea căilor alternative în mod automat la configurarea traseului principal.

După cum se observă în figura 5.3, folosirea protecției locale pentru fiecare LSR din traseul principal pastrează multe resurse pentru scopuri de rezervă. Deși există doar un singur traseu principal în rețeaua din exemplu, există unele legături pe care se fac doua rezervări pentru backup. Acest lucru nu este necesar deoarece nu există decât un singur traseu principal și acele rezervări nu pot fi folosite în același timp dacă are loc un singur eșec pe traseul principal. Acest lucru este legat de rezervările LSR7 – LSR8 și LSR9 – LSR5, și nu de rezervarea dintre LSR3 și LSR8 deoarece avem doua direcții opuse. Pentru a preveni această rezervare dublă, se pot unifica acele rezervări și pot fi partajate. Unificarea se poate realiza dacă cele două sau mai multe rezervări sunt rezervări de backup pentru același traseu principal. Sunt folosite următoarele reguli pentru unificarea rezervărilor:

Pentru toate rezervările care partajează aceeași interfața de ieșire și același LSR next-hop, rezervarea se poate efectua doar pe traseul cel mai scurt către nodul de ieșire al LSP-ului protejat. Dacă traseele au același număr de hop-uri, traseul către străbate noduri pe care altele le evită, atunci acel traseu ar trebui eliminat. Dacă încă mai există trasee alternative multiple, atunci este o decizie locală care se va alege.

În exemplul de mai sus, acest lucru înseamnă că pe legătura dintre LSR7 și LSR8, rezerva LSR1 va fi unită cu rezerva LSR2 în LSR7. Acest lucru se întâmplă deoarece cea de-a doua rezervă are un traseu mai scurt către nodul de ieșire. LSR9 va unifica rezervele LSR3 și LSR4 pe legătura dintre LSR8 și LSR9. Alternativa de la LSR3 este proiectată să ocolească LSR4, așa că alternativa de la LSR3 va face ultima rezervare.

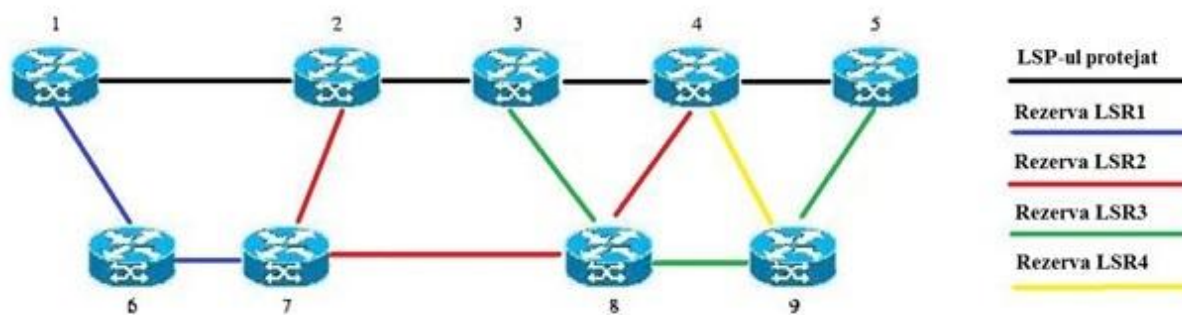


Figura 5.4: Redirijarea rapidă cu rezervă unu la unu și unificare

Traseul de la LSR7 la LSR4 este folosit acum de rezerva lui LSR1 și rezerva LSR2 iar traseul de la LSR9 la LSR5 este folosit de rezerva lui LSR3 și rezerva lui LSR4. Folosind această tehnică de unificare, putem observa că de la 12 rezervări inițiale am ajuns la 9 rezervări inițiale.

5.4.2 Redirijarea rapidă cu rezervă de facilitate (Fast reroute Facility Backup)

Folosind acest model, este creat un singur LSP care este folosit ca alternativă la multiple LSP-uri în loc să folosească o rezervă separată pentru fiecare LSP protejat. Acest lucru constrânge setul de LSP-uri protejate să traverseze un LSR comun pentru a putea fi redirijate pe calea de rezervă. Toate LSP-urile care traversează un anumit punct de reparare și acest LSR comun pot fi protejate de tunelul ocolitor. Rezultă că modelul cu rezervă de facilitate folosește asocieri ale căilor unu la unu.

Tunelul folosește suprapunerea etichetelor, deoarece, atunci când mai multe LSP-uri sunt protejate de un singur tunel de ocolire, trebuie să existe o cale să separe LSP-urile diferite când traficul de pe tunelul de rezervă ajunge la nodul de ieșire.

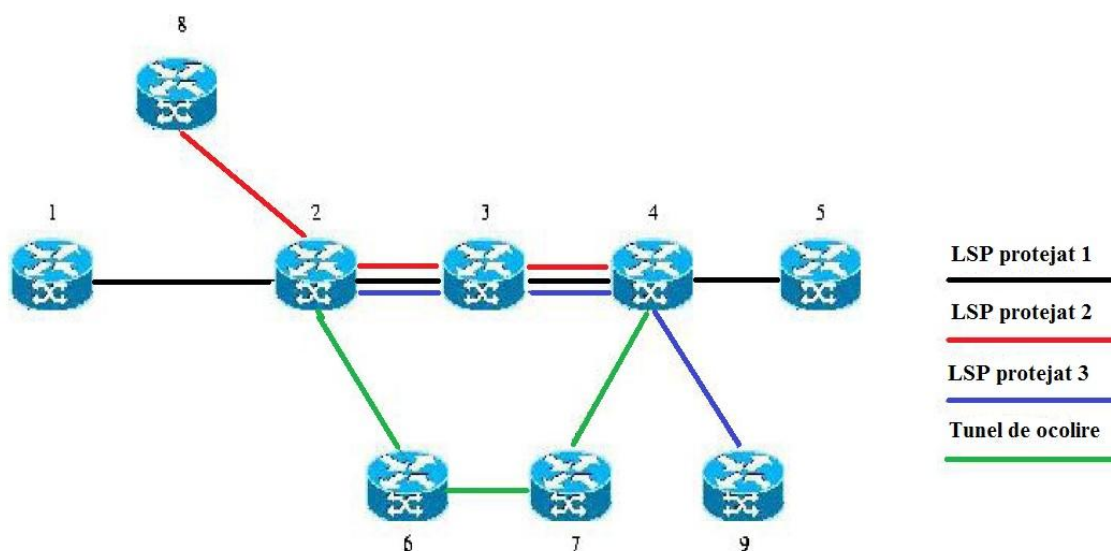


Figura 5.5: Modelul de redirijare rapidă cu rezervă de facilitate

În figura de mai sus avem 3 LSP-uri care sunt protejate de un singur tunel de ocolire. După cum se observă, tunelul de rezervă cu facilitate folosește restaurarea locală cu protecția nodului. În acest caz, LSR3 este nodul protejat pentru cele trei LSP-uri. Dacă s-ar fi folosit tehnica unu la unu, s-ar fi utilizat trei legături de rezervă pentru a proteja LSR3. Chiar dacă mai multe LSP-uri pot partaja același tunel de ocolire, resursele nu sunt partajate. Fiecare LSP protejat trebuie să facă o rezervare adițională dacă vrea să folosească tunelul.

5.5 Soluția propusă de restabilire a traficului

Având în vedere modelele anterioare de restabilire a traficului, am propus o soluție care combină comutarea traficului, specifică redirijării rapide, cu stabilirea unui traseu disjunct total față de traseul principal, specific metodelor Makam și Haskin. Această soluție caută să îmbunătățească numărul de resurse folosite în modelul redirijării rapide.

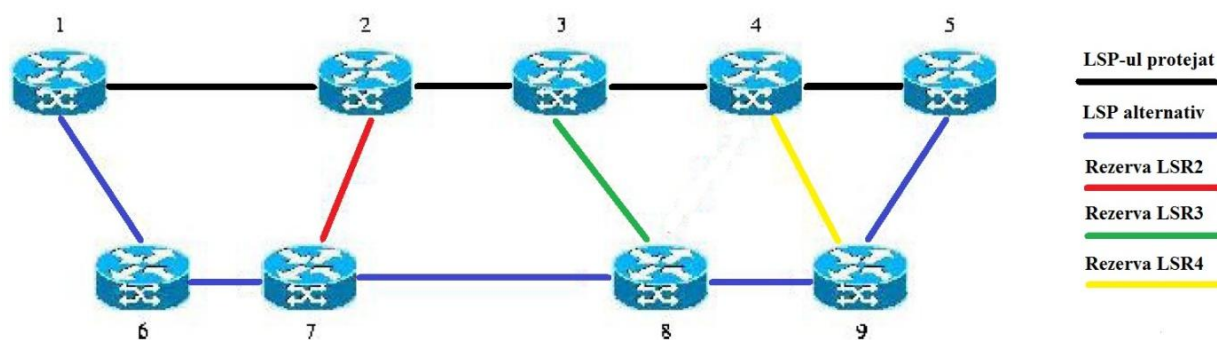


Figura 5.6: Soluția propusă de stabilire a LSP-urilor de rezervă

Această schemă de restabilire presupune formarea LSP-urilor de rezervă înaintea eșecului. Primul LSP creat trebuie să fie un traseu disjunct față de calea protejată care va uni nodul de intrare în rețeaua MPLS cu nodul de ieșire. Acesta va fi și traseul folosit în cazul în care LSR-ul de intrare va fi cel care detectează eșecul în LSP-ul protejat.

Urmatoarele LSP-uri de rezervă vor fi create astfel încât pentru fiecare LSR din traseul principal se va calcula drumul cel mai scurt către LSP-ul alternativ și disjunct. Pentru a găsi cel mai scurt drum se va folosi nivelul de rețea deoarece etichetele MPLS nu sunt suficiente, ele oferind acces doar la LSR-urile vecine. Deși calcularea acestor traseuri va implica schimburi de mesaje IP între LSR-uri, acest lucru se face doar o singură dată, când se creează toate LSP-urile de rezervă.

Funcționarea acestui model este următoarea: atunci când un LSR din calea protejată detectează un eșec, acesta va comuta traficul pe traseul său de rezervă. Atunci când acest trafic ajunge la un LSR prin care trece primul LSP de rezervă creat (LSP-ul alternativ), va fi comutat pe acest LSP către nodul de ieșire.

Capitolul 6

6. Simulări

6.1 Simulatorul NS2

NS2 este un simulator în care un nod este format din clasificatori și agenți. Un agent este obiectul care primește sau recepționează, în timp ce clasificatorul este un obiect responsabil cu clasificarea pachetelor și cu dirijarea lor către nodurile convenabile sau cu livrarea către agentul local dacă nodul receptor este destinația pachetului. Astfel, pentru a construi un nod MPLS, un clasificator nou numit „Clasificator MPLS” ar trebui creat pentru a putea clasifica pachetele recepționate, a determina dacă sunt etichetate sau nu și tratate în mod corespunzător. De asemenea, un nou agent, „Agent LDP” trebuie inserat în nodul IP pentru a distribui etichetele între noduri MPLS și pentru a construi LSP-urile.

Arhitectura unui nod MPLS incluzând clasificatorul MPLS și agentul LDP:

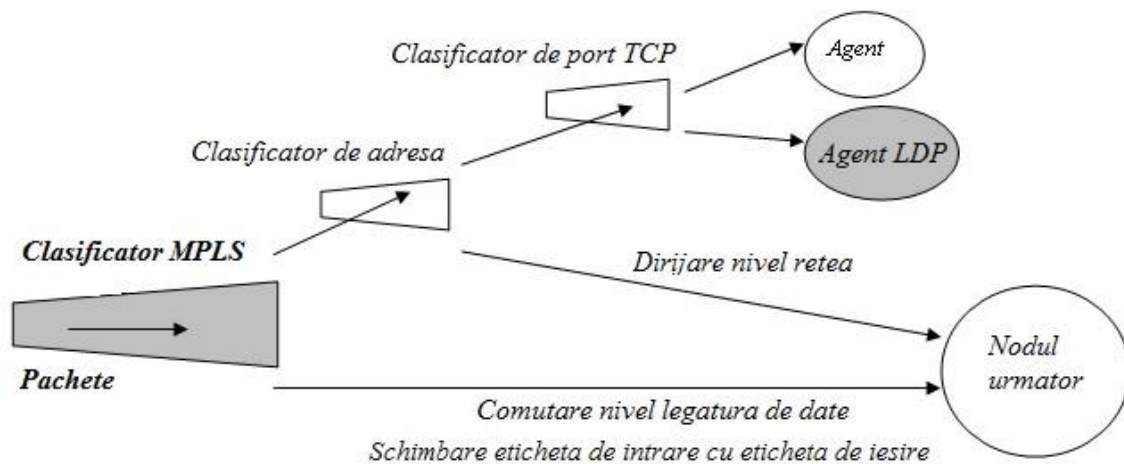


Figura 6.1: Arhitectura unui nod MPLS în NS2

Atunci când un nod recepționează un pachet, clasificatorul MPLS determină dacă este etichetat sau nu. Dacă pachetul este etichetat, clasificatorul va executa o comutare la nivel de legătură de date. Va înlocui eticheta prezentă în antetul MPLS al pachetului cu eticheta corespunzătoare destinației pachetului (FEC) după care îl va trimite către nodul următor. Dacă pachetul recepționat nu este etichetat și există un LSP creat pentru destinația pachetului (LSP) atunci clasificatorul va crea un antet MPLS, îi va asocia o etichetă și îl va trimite către next-hop-ul corespunzător.

În cazul în care pachetul recepționat nu este etichetat și nu există LSP, nodul va livra pachetul clasificatorului de adresă care va executa dirijare obișnuită de nivel rețea prin examinarea adresei IP destinație.

Pachetul va fi livrat clasificatorului de port când nodul receptor este și destinația pachetului. Acest clasificator de port va furniza pachetul agentului convenabil. Agentul LDP este folosit pentru a distribui etichetele și pentru a inițializa LSP-uri pe baza protocolului LDP.

Un nod MPLS are trei tabele pentru a organiza informația legată de LSP-uri și de distribuția tabelelor. Cele trei tabele sunt:

- Tabel cu informații despre etichete(Label Information Base – LIB): Acest tabel este construit și folosit pentru a asocia eticheta și interfața de intrare cu eticheta și interfața de ieșire. Este folosit atunci când se efectuează comutarea la nivel de legătură de date.
- Tabel cu comutări parțiale(Partial Forwarding Table - PFT): Acest tabel este folosit atunci când sosește un pachet neetichetat. Nodul MPLS va cauta în acest tabel pentru a găsi FEC-ul reprezentat de adresa destinație a pachetului. Instanța găsită ar putea indica o intrare în tabelul LIB pentru a eticheta pachetul sau ar putea indica NULL și ca rezultat se va recurge la dirijare obișnuită la nivel de rețea.
- Tabel cu informații de rutare explicite(Explicit Routing information Base - ERB): Acest tabel este folosit pentru a stoca informații cu rute explicite. Dacă se dorește o rută explicită, singurul lucru de făcut este inserarea unei instanțe care va conține un indicator către același pointer LIB din tabelul PFT.

Network Animator(Nam)

Nam este o unealtă pentru animații, bazată pe Tcl/Tk , folosită pentru vizualizarea datelor simulărilor de rețele. Teoria de implementare din spatele Nam a fost crearea unui animator care să poată citi seturi foarte mari de date pentru animații și care să poată fi îndeajuns de extensibil pentru a putea vizualiza o varietate foarte mare de situații de rețele. Având această constrângere, nam a fost creat să poată citi comenzi simple de animație din interiorul unor fișiere foarte mari de urmărire. Comenzile sunt reținute în fișier și sunt recitite ori de câte ori este necesar.

Primul pas pentru a putea folosi nam este crearea unui fișier de urmărire. Acest fișier conține informații referitoare la topologii, noduri, legături dar și pachete. De obicei, acest fișier este creat de către Network Simulator.

După ce a fost generat fișierul pentru urmărire, acesta este pregătit pentru a putea fi animat de către nam. La pornire, nam citește fișierul, crează topologia, creează o fereastră pop-up, realizează layout-ul necesar și se oprește la momentul de timp 0. Network animator are o interfață pentru utilizator însă permite controlul asupra multor aspecte ale animației.

6.2 Scenariul simulat

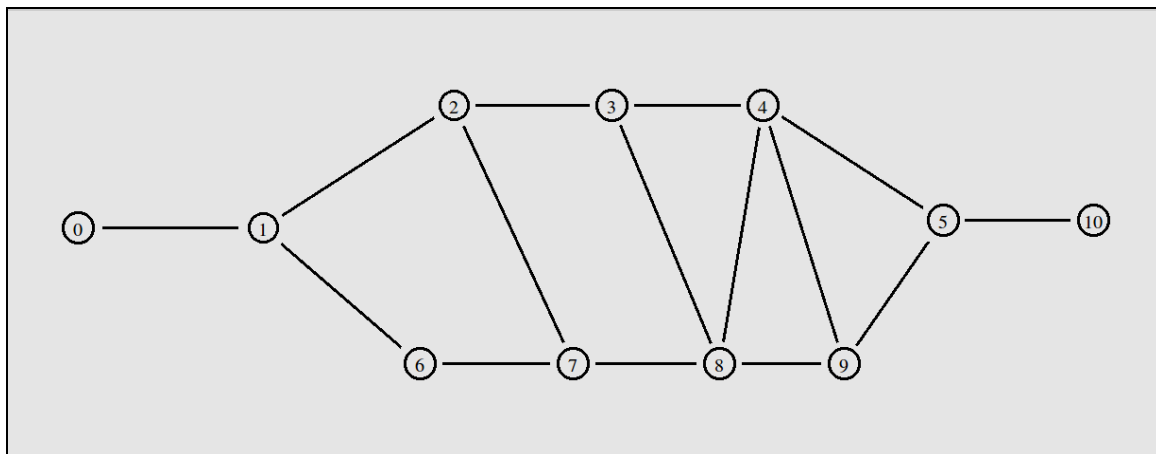


Figura 6.2: Topologia utilizată

Topologia utilizată este alcătuită din 9 noduri MPLS și 2 noduri IP(0 și 10).

Considerăm că nodul 0 este sursa traficului și nodul 10 este destinația. Inițial, traficul se va trimite folosind traseul protejat(1-2-3-4-5) și se vor folosi alte trasee atunci când rețeaua va suferi un eșec.

În simulare am folosit eșecul unei singure legături de pe traseul principal, legătura dintre LSR4 și LSR5. Timpul de detecție al eșecului este 0 pentru a nu avea diferențe între modelele folosite.

Toate legăturile între noduri sunt identice, legături de 100Mb și întârziere de 10ms. Pentru a varia rata de transfer, am folosit diferite dimensiuni ale pachetelor, începând cu 100 octeți până la 10000 octeți.

Numărul de pachetele pierdute se calculează ca diferența între numărul de pachete trimise de nodul sursa(0) și numărul de pachete recepționate de nodul destinație(10).

Tabelul 6.1 Date legate de scenariul simulat

Tipul cozii de așteptare	DropTail
Numărul de noduri IP	2
Numărul de noduri MPLS	9
Durata timpului simulat	1s
Timpul de detecție al eșecului	0s
Lărgimea de bandă a legăturilor	100 MB/s
Tipul de trafic	UDP
Dimensiunea pachetelor	100/500/1000/5000/10000 bytes
Intervalul de transmitere al Pachetelor	0,001s
Protocolul de distribuție al etichetelor	LDP

Pentru realizarea scenariilor am creat scripturi în limbajul TCL pentru fiecare model de recuperare de trafic. Acestea se găsesc în anexa.

Pentru a extrage datele din fișierul rezultat în urma simulării(trace file) am folosit scripturi realizate în limbajul AWK. Sunt folosite două scripturi, unul care calculează timpul de întrerupere al serviciului și unul care calculează numărul de pachete pierdute.

Capitolul 7

7. Prezentarea rezultatelor

Tabel 7.1 – Rezultatele obținute pentru diverse metode de recuperare

Metoda de recuperare	Timp de întrerupere al serviciului (secunde)	Pachete pierdute	Resurse rezervate
Makam	0,0858	74	5
Haskin	0,0858	11	8
Redirijare	0,0648	53	0
Redirijare rapidă	0,0228	11	12
Propunere	0,0228	11	8

Tabel 7.2 – Numărul de pachete pierdute în funcție de rata de transfer

Metoda de recuperare	Numărul de pachete pierdute				
	Rata de transfer				
	0,1 MB/s	0,5 MB/s	1MB/s	5 MB/s	10 MB/s
Makam	74	70	70	350	702
Haskin	11	10	10	50	504
Redirijare	53	51	51	251	530
Redirijare rapidă	11	10	10	50	101
Propunere	11	10	10	50	101

Prin varierea ratei de transfer nu se afectează numărul de resurse rezervate și nici timpul de întrerupere al serviciului nu are decât variații foarte mici și de aceea nu ar fi avut sens prezentarea acestor date. Pentru rate de transfer de 0,1 MB/s și 0,5 MB/s nu avem diferențe mari între numărul de pachete pierdute deoarece nu se variază decât dimensiunea acestora și nu numărul lor.

Pentru verificarea rezultatelor am folosit și o rețea cu doar 7 noduri MPLS, dar rezultatele obținute au fost asemănătoare și prezentarea lor nu ar fi avut sens.

7.1 Timpul de întrerupere al serviciului folosind diverse metode de recuperare

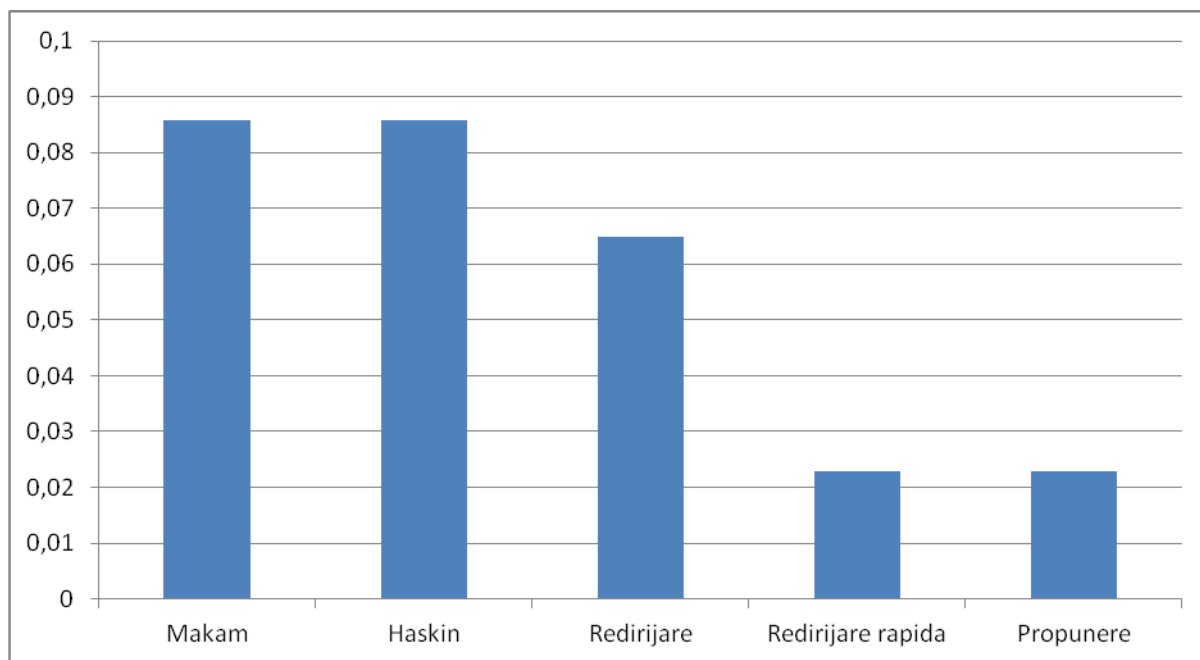


Figura 7.1: Timpul de întrerupere al serviciului metodelor de recuperare

Timpul de întrerupere al serviciului se măsoară ca timpul trecut de la recepționarea ultimului pachet în nodul 10, ca rezultat al căderii legăturii, până când sosește primul pachet după ce a fost restabilit traficul.

După cum se observă din figura de mai sus, metodele Makam, Haskin și metoda redirijării nu sunt potrivite pentru transferul de voce deoarece au timpul de întrerupere al serviciului mai mare de 60ms. Orice întârziere mai mare de 60ms în semnalul de voce este percepută ca o pauză în vorbire și este întrerupt fluxul vorbirii.

Pentru metoda redirijării rapide și pentru soluția propusă, timpii de întrerupere al serviciului sunt minimi deoarece comutarea se face local pe traseele de rezervă stabilite anterior. Oriunde ar avea loc un eșec de-a lungul LSP-ului protejat, traficul este comutat pe rezerva corespunzătoare și astfel timpul de întrerupere al serviciului va fi constant.

Metodele Makam și Haskin au performanțe asemănătoare deoarece pentru a redirija traficul trebuie notificat nodul de intrare despre eșecul rețelei. Cu cât eșecul se produce mai departe de nodul de intrare, cu atât va fi mai mare timpul de întrerupere al serviciului.

7.2 Numărul de pachete pierdute folosind diverse metode de recuperare

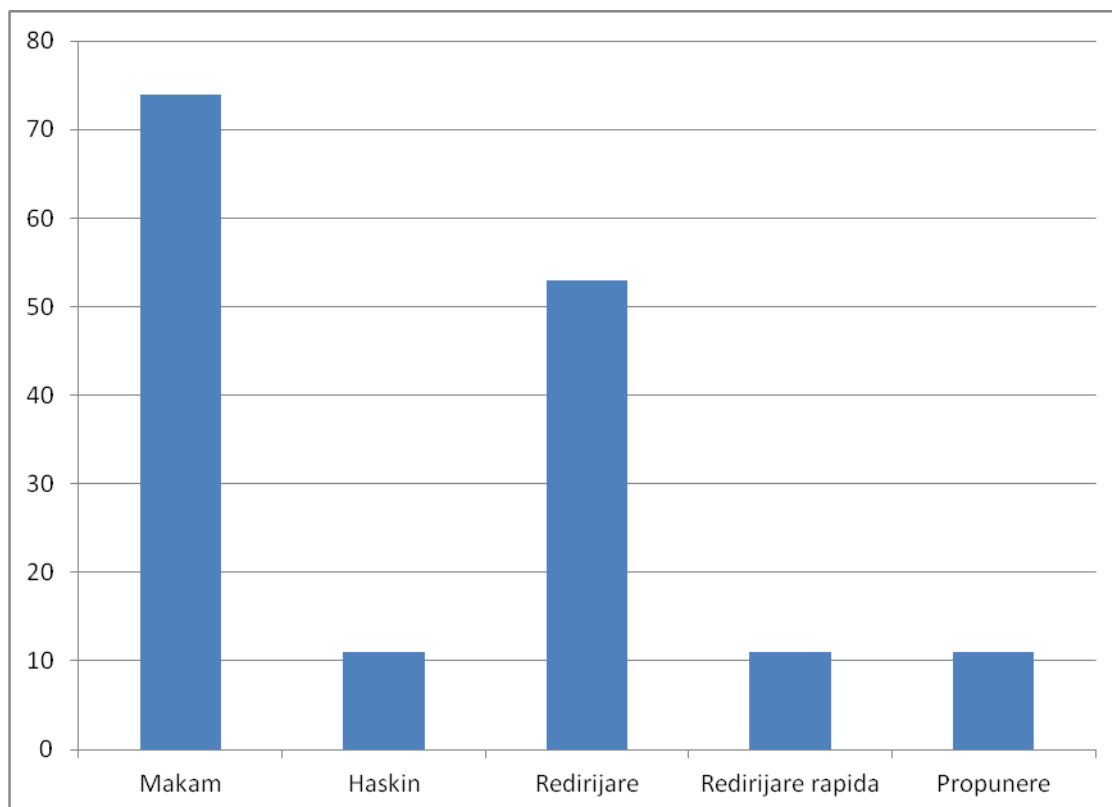


Figura 7.2: Numărul de pachete pierdute folosind diverse metode de recuperare

Numărul de pachetele pierdute se calculează ca diferența între numărul de pachete trimise de nodul sursă și numărul de pachete recepționate de nodul destinație.

Se observă în figura de mai sus că se pierde multe pachete folosind metoda Makam deoarece pachetele sunt trimise pe un traseu eșuat până când nu este notificat nodul de intrare să comute acest trafic pe calea de rezervă. Numărul de pachete pierdute va scădea cu cât eșecul se produce mai aproape de nodul de intrare.

Metoda redirijării trebuie să găsească un traseu alternativ după ce a fost detectat un eșec și astfel pachetele se vor pierde până când nu se va găsi un traseu potrivit.

Numărul de pachete pierdute este minim în cazul metodelor Haskin, redirijării rapide și soluției propuse deoarece traficul este comutat pe căile de rezervă în momentul în care a fost detectat un eșec. Folosind metoda Haskin, ar putea fi probleme cu protocoale de nivel superior deoarece pachetele nu vor sosi în ordinea în care au fost trimise.

7.3 Numărul de resurse rezervate folosind diverse metode de recuperare

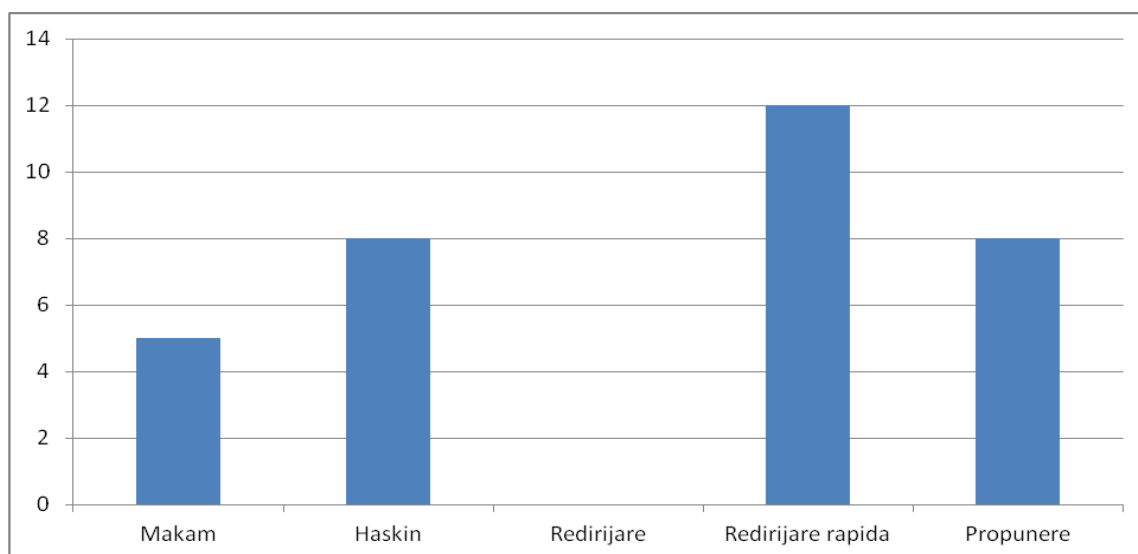


Figura 7.3: Numărul de resurse rezervate folosind diverse metode de recuperare

Numărul de resurse rezervate se determină ca fiind numărul total de legături din toate LSP-urile create și nu numărul de LSP-uri create.

În exemplul folosit, modelul Makam folosește 5 rezervări de resurse. Modelul Haskin va folosi întotdeauna mai multe deoarece, în plus față de traseul global de rezervă, va avea nevoie de un traseu pentru traficul inversat. Modelul redirijării nu folosește deloc resurse rezervate anterior, deoarece calea de rezervă se va crea când va fi nevoie, la producerea unui eșec.

Numărul de LSP-uri folosit de metodele redirijării rapide și soluției propuse este identic și egal cu $N-1$, N fiind numărul de noduri din LSP-ul protejat. Deoarece metoda redirijării rapide caută să găsească un LSP de rezervă pentru fiecare nod din traseul principal și apoi să revină la traseul principal se folosesc mai multe resurse față de soluția propusă.

În acest exemplu, numărul de resurse rezervate este același pentru modelul Haskin și pentru soluția propusă deoarece avem o distanță de un singur hop între nodurile de pe traseul principal și nodurile de pe LSP-ul alternativ. Într-o rețea de mari dimensiuni, numărul de resurse rezervate va fi mai mare pentru soluția propusă decât metoda Haskin și mult mai mare pentru metoda redirijării rapide.

Numărul de resurse rezervate depinde de topologia rețelei. Deși metoda redirijării rapide cu rezervă unu la unu și unificare folosește mai puține resurse, acesta pornește de la premisa ca nu putem avea mai mult de un singur eșec la un moment dat.

7.4 Numărul de pachete pierdute în funcție de rata de transfer

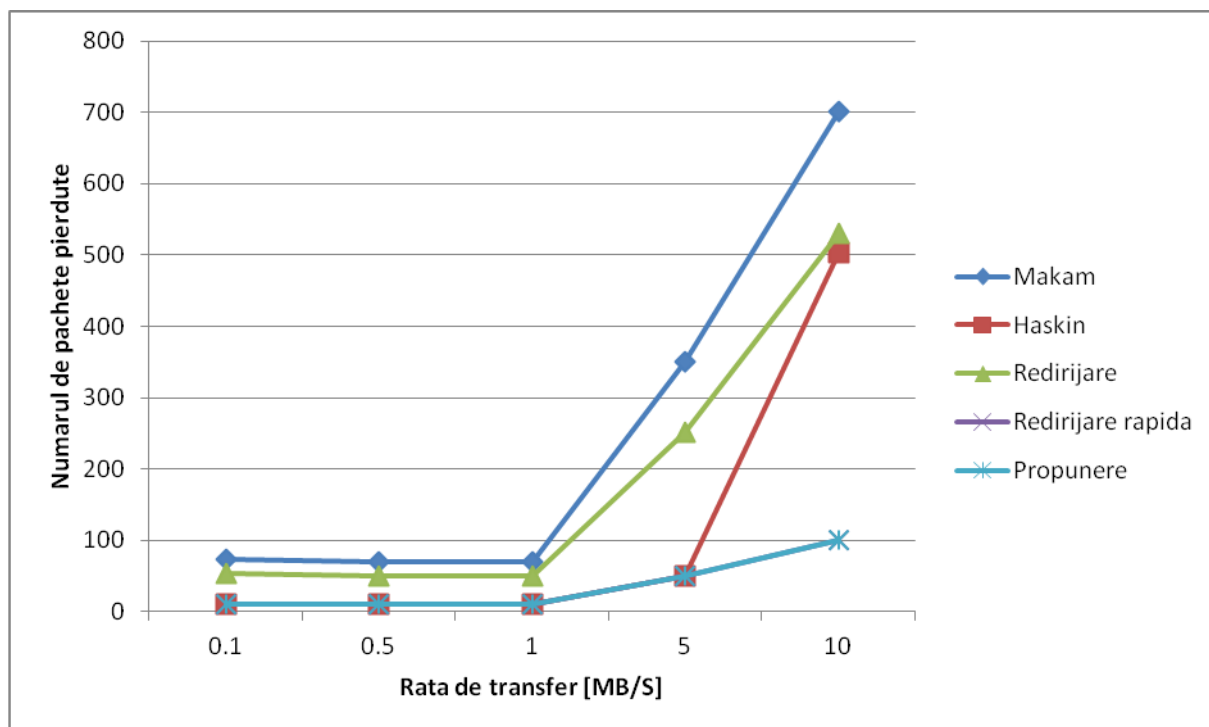


Figura 7.4: Numărul de pachete pierdute în funcție de rata de transfer

În figura de mai sus, pe axa orizontală avem dimensiunea pachetelor și pe axa verticală avem numărul de pachete pierdute. După cum se observă, cu cât crește rata de transfer, cu atât vor fi mai multe pachete pierdute.

Pentru metodele redirijării rapide și soluției propuse avem aceleași valori de pachete pierdute și graficele se suprapun. Defapt singurele pachete care se pierd la aceste două metode sunt pachetele care se află pe legătură în momentul când aceasta eșuează deoarece timpul de detecție al eșecului este 0.

Metoda Makam are pierderi foarte mari de pachete și nu este potrivită pentru rețele în care rata de transfer este foarte mare. Cam același lucru se poate spune și despre metoda redirijării, deși pierderile sunt mai mici.

Metoda Haskin trebuie să inverseze traficul spre nodul de intrare iar la sosire este amestecat cu traficul de pe LSP-ul alternativ și se depășește lățimea de bandă pentru acest LSP alternativ. Astfel apare congestia traficului și, în exemplul meu, pachetele au fost aruncate deoarece nu am folosit nici un mecanism de tratare a congestiei. Din această cauză se observă în figură că pentru rate mari de transfer, se pierd mai multe pachete cu metoda Haskin decât cu metoda redirijării rapide sau metoda propusă.

Concluzii

Tehnologia MPLS a combinat capabilitățile de comutare a etichetei și noțiunea de circuit virtual specifice ATM cu orientarea pachetelor, specifică tehnologiei IP. Principalul avantaj MPLS îl reprezintă flexibilitatea în sensul că a fost proiectat să lucreze într-un mediu cu protocoale multiple.

Într-o rețea WAN modernă, trebuie ca utilizatorii să nu aibă de suferit în cazul în care are loc o cădere a unui nod sau a unei legături. În mod ideal ar trebui ca o rețea să aibă redundanța totală, dar practic acest lucru nu este întotdeauna posibil din cauza costului. Pentru a studia o rețea fără a fi nevoie de implementarea ei fizică, am studiat comparativ diverse simulatoare de rețea și pot afirma că sunt niște instrumente foarte utile în domeniul rețelisticii.

În această lucrare am analizat diverse metode de recuperare a traficului în rețele MPLS și am propus o soluție pentru îmbunătățirea numărului de resurse folosite în metoda redirejării rapide fără a sacrifica capabilitatea de a trata mai multe eșecuri la un moment dat. În topologia utilizată, s-a realizat o îmbunătățire de 33,3% dar în alte topologii poate varia foarte mult.

În final, pot afirma că nu există o metodă de recuperare a traficului care să fie cea mai bună în mod exclusiv, fiecare metodă are avantajele și dezavantajele sale și această lucrare ar putea ajuta administratorii de rețele să facă o alegere bună în această privință.

Bibliografie și referințe

- [1]- Andrew S. Tanenbaum - Computer networks, Prentice Hall Professional Technical Reference, 4th edition, 2002, pag. 37-39, 75-78, 392-394;
- [2]- Johan Martin Olof Petersson - MPLS Based Recovery Mechanics, University of Oslo, May 2005;
- [3]- Lemma Hundessa Gonfa - Enhanced Fast Rerouting Mechanisms for Protected Traffic in MPLS Networks, University of Catalunya, February 2003;
- [4]- Dr. Yogesh P Kosta, Minesh P Thaker, Bhaumik Nagar - Performance Comparison of Rerouting Schemes of Multi Protocol Label Switching Network, ACEEE Int. J. on Control System and Instrumentation, Vol. 02, No. 01, Feb 2011;
- [5]- Gaeil Ahn, Jongsoo Jang, Woojik Chun - An Efficient Rerouting Scheme for MPLS-Based Recovery and Its Performance Evaluation, Telecommunication Systems 19:3,4, 481–495, 2002;
- [6]- Ravi Tomar, Anuj Kumar Yadav, Deep Kumar, Prashant Chaudhary - MPLS-A Fail Safe Approach to Congestion Control, International Journal of Advanced Research in Computer Science and Software Engineering, Volume 2, Issue 7, July 2012;
- [7]- Ed Harrison, Ben Miller, Adrian Farrel - Protection And Restoration In MPLS Networks, An examination of the methods for protecting MPLS LSPs against failures of network resources, First issued October 2001;
- [8]- K. Owens, V. Sharma, S. Makam, and C. Huang - A Path Protection Restoration Mechanism for MPLS Networks. Work in progress, Internet draft <draft-chang-mpls-protection-03.txt>, July 2001;
- [9]- D. Haskin and R. Krishnan - A Method for Setting an Alternative Label Switched Paths to Handle Fast Reroute. Work in progress, Internet draft <draft-haskin-mpls-fast-reroute-05.txt>, November 2000;
- [10]- P. Pan, G. Swallow, A. Atlas - “Fast Reroute Extensions to RSVP-TE for LSP Tunnels” Internet Draft draft-ietf-mpls-rsvp-lsp-fastreroute-06.txt, June 2004;
- [11]- Ginsburg, 1996; Goralski, 1995; Ibe, 1997; Kim et al., 1994; at Stallings, 2000;
- [12]- Cemal Kocak, Ismail Erturk, Huseyin Ekiz - Comparative Performance Analysis of MPLS over ATM and IP over ATM Methods for Multimedia Transfer Applications;

- [13]- Gaeil Ahn, Woojik Chun – Overview of MPLS Network Simulator: Design and Implementation, Department of Computer Engineering, Chungnam National University, Korea;
- [14]- John Cleary, Murray Pearson, Ian Graham, Brian Unger - High Performance Simulation for ATM Network Development, Department of Computer Science, University of Waikato, June 1996;
- [15]- Mrs. Saba Siraj, Mr. Ajay Kumar Gupta, Mrs Rinku-Badgujar - Network Simulation Tools Survey, International Journal of Advanced Research in Computer and Communication Engineering, Vol. 1, Issue 4, June 2012;
- [16]- Paul Meenaghan and Declan Delaney - An Introduction to NS, Nam and OTcl Scripting, National University of Ireland, Maynooth, Maynooth, Co. Kildare, Ireland, Department of Computer Science, 2004;
- [17]- Elias Weingartner, Hendrik vom Lehn and Klaus Wehrle - A performance comparison of recent network simulators, Distributed Systems Group, RWTH Aachen University, Aachen, Germany;
- [18]- Jianli Pan - A Survey of Network Simulation Tools: Current Status and Future Developments, <http://www.cse.wustl.edu/~jain/cse567-08/index.html>;
- [19]- http://en.wikipedia.org/wiki/WAN_optimization;
- [20]- <http://searchenterprisewan.techtarget.com/feature/WAN-optimization-techniques>;
- [21]- <http://www.bcs.org/content/conwebdoc/6849>;
- [22]- http://en.verysource.com/code/2773401_1/test-reroute.tcl.html;

Anexe

Makam.tcl

```
set ns [new Simulator]

set nf [open MPLSexample.nam w]
$ns namtrace-all $nf
set f0 [open MPLSexample.tr w]
$ns trace-all $f0

proc finish {} {
    global ns nf f0
    $ns flush-trace
    close $nf
    close $f0
    exec nam MPLSexample.nam &
    exit 0
}

#Agent/LDP set trace_ldp_ 1
#Classifier/Addr/MPLS set trace_mpls_ 1

#$ns rtproto DV

set node0 [$ns node]
$ns node-config -MPLS ON
set LSR1 [$ns node]
set LSR2 [$ns node]
set LSR3 [$ns node]
set LSR4 [$ns node]
set LSR5 [$ns node]
set LSR6 [$ns node]
set LSR7 [$ns node]
set LSR8 [$ns node]
set LSR9 [$ns node]
$ns node-config -MPLS OFF
set node10 [$ns node]

$ns duplex-link $node0 $LSR1 100Mb 10ms DropTail
$ns duplex-link $LSR1 $LSR2 100Mb 10ms DropTail
$ns duplex-link $LSR2 $LSR3 100Mb 10ms DropTail
$ns duplex-link $LSR3 $LSR4 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR5 100Mb 10ms DropTail
$ns duplex-link $LSR5 $node10 100Mb 10ms DropTail
$ns duplex-link $LSR1 $LSR6 100Mb 10ms DropTail
$ns duplex-link $LSR6 $LSR7 100Mb 10ms DropTail
$ns duplex-link $LSR7 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR8 $LSR9 100Mb 10ms DropTail
$ns duplex-link $LSR9 $LSR5 100Mb 10ms DropTail
$ns duplex-link $LSR2 $LSR7 100Mb 10ms DropTail
$ns duplex-link $LSR3 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR9 100Mb 10ms DropTail

#
# configure ldp agents on all mpls nodes
#
for {set i 1} {$i < 10} {incr i} {
```

```

    set a LSR$i
    for {set j [expr $i+1]} {$j < 10} {incr j} {
        set b LSR$j
        eval $ns LDP-peer $$a $$b
    }
    set m [eval $$a get-module "MPLS"]
    $m enable-reroute "new"
}

#configuram culori pentru diferitele pachete LDP
$ns ldp-request-color      blue
$ns ldp-mapping-color      red
$ns ldp-withdraw-color     magenta
$ns ldp-release-color      orange
$ns ldp-notification-color yellow

Classifier/Addr/MPLS enable-on-demand
Classifier/Addr/MPLS enable-ordered-control

set src0 [new Application/Traffic/CBR]
set udp0 [new Agent/UDP]
$src0 attach-agent $udp0
$ns attach-agent $node0 $udp0
$src0 set packetSize_ 10000
$src0 set interval_ 0.001

set Dst0 [new Agent/LossMonitor]
$ns attach-agent $node10 $Dst0

$ns connect $udp0 $Dst0

$ns at 0.1 "$src0 start"

#protected LSP
$ns at 0.1 "[$LSR1 get-module MPLS] make-explicit-route 5 2_3_4_5 1000 -1"
$ns at 0.1 "[$LSR1 get-module MPLS] flow-erlsp-install 10 -1 1000"

#alternate LSP
$ns at 0.1 "[$LSR1 get-module MPLS] make-explicit-route 5 6_7_8_9_5 1001 -1"

#link failure
$ns rtmodel-at 0.3 down $LSR4 $LSR5
$ns at 0.3 "[$LSR4 get-module MPLS] make-explicit-route 1 3_2_1 1002 -1"
$ns at 0.3 "[$LSR4 get-module MPLS] ldp-trigger-by-withdraw 10 -1"

#change flow to backup LSP
$ns at 0.33 "[$LSR1 get-module MPLS] flow-erlsp-install 10 -1 1001"

#restore link
$ns rtmodel-at 0.5 up $LSR4 $LSR5

$ns at 0.6 "$src0 stop"

$ns at 1.0 "finish"
$ns run

```

Haskin.tcl

```
set ns [new Simulator]

set nf [open MPLSexample.nam w]
$ns namtrace-all $nf
set f0 [open MPLSexample.tr w]
$ns trace-all $f0

proc finish {} {
    global ns nf f0
    $ns flush-trace
    close $nf
    close $f0
    exec nam MPLSexample.nam &
    exit 0
}

#Agent/LDP set trace_ldp_ 1
#Classifier/Addr/MPLS set trace_mpls_ 1

$ns rtproto DV

set node0 [$ns node]
$ns node-config -MPLS ON
set LSR1 [$ns node]
set LSR2 [$ns node]
set LSR3 [$ns node]
set LSR4 [$ns node]
set LSR5 [$ns node]
set LSR6 [$ns node]
set LSR7 [$ns node]
set LSR8 [$ns node]
set LSR9 [$ns node]
$ns node-config -MPLS OFF
set node10 [$ns node]

$ns duplex-link $node0 $LSR1 100Mb 10ms DropTail
$ns duplex-link $LSR1 $LSR2 100Mb 10ms DropTail
$ns duplex-link $LSR2 $LSR3 100Mb 10ms DropTail
$ns duplex-link $LSR3 $LSR4 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR5 100Mb 10ms DropTail
$ns duplex-link $LSR5 $node10 100Mb 10ms DropTail
$ns duplex-link $LSR1 $LSR6 100Mb 10ms DropTail
$ns duplex-link $LSR6 $LSR7 100Mb 10ms DropTail
$ns duplex-link $LSR7 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR8 $LSR9 100Mb 10ms DropTail
$ns duplex-link $LSR9 $LSR5 100Mb 10ms DropTail
$ns duplex-link $LSR2 $LSR7 100Mb 10ms DropTail
$ns duplex-link $LSR3 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR9 100Mb 10ms DropTail

#
# configure ldp agents on all mpls nodes
#
for {set i 1} {$i < 10} {incr i} {
    set a LSR$i
    for {set j [expr $i+1]} {$j < 10} {incr j} {
        set b LSR$j
        eval $ns LDP-peer $$a $$b
    }
}
```

```

    }
    set m [eval $$a get-module "MPLS"]
    $m enable-reroute "new"
}

#configuram culori pentru diferitele pachete LDP
$ns ldp-request-color      blue
$ns ldp-mapping-color      red
$ns ldp-withdraw-color     magenta
$ns ldp-release-color      orange
$ns ldp-notification-color yellow

Classifier/Addr/MPLS enable-on-demand
Classifier/Addr/MPLS enable-ordered-control

set src0 [new Application/Traffic/CBR]
set udp0 [new Agent/UDP]
$src0 attach-agent $udp0
$ns attach-agent $node0 $udp0
$src0 set packetSize_ 10000
$src0 set interval_ 0.001

set Dst0 [new Agent/LossMonitor]
$ns attach-agent $node10 $Dst0

$ns connect $udp0 $Dst0

$ns at 0.1 "$src0 start"

#alternative LSP
$ns at 0.1 "[$LSR1 get-module MPLS] make-explicit-route 5 6_7_8_9_5 1000 -1"

#backward LSP
$ns at 0.2 "[$LSR5 get-module MPLS] make-explicit-route 5 4_3_2_1_1000 1005 -1"

#link failure
$ns rtmodel-at 0.3 down $LSR4 $LSR5

#change the working LSP to backward LSP
$ns at 0.3 "[$LSR4 get-module MPLS] flow-erlsp-install 10 -1 1005"
$ns at 0.33 "[$LSR1 get-module MPLS] flow-erlsp-install 10 -1 1000"

#restore link
$ns rtmodel-at 0.5 up $LSR4 $LSR5

$ns at 0.6 "$src0 stop"

$ns at 1.0 "finish"
$ns run

```

Reroute.tcl

```
set ns [new Simulator]

set nf [open MPLSexample.nam w]
$ns namtrace-all $nf
set f0 [open MPLSexample.tr w]
$ns trace-all $f0

proc finish {} {
    global ns nf f0
    $ns flush-trace
    close $nf
    close $f0
    exec nam MPLSexample.nam &
    exit 0
}

#Agent/LDP set trace_ldp_ 1
#Classifier/Addr/MPLS set trace_mpls_ 1

$ns rtproto DV

set node0 [$ns node]
$ns node-config -MPLS ON
set LSR1 [$ns node]
set LSR2 [$ns node]
set LSR3 [$ns node]
set LSR4 [$ns node]
set LSR5 [$ns node]
set LSR6 [$ns node]
set LSR7 [$ns node]
set LSR8 [$ns node]
set LSR9 [$ns node]
$ns node-config -MPLS OFF
set node10 [$ns node]

$ns duplex-link $node0 $LSR1 100Mb 10ms DropTail
$ns duplex-link $LSR1 $LSR2 100Mb 10ms DropTail
$ns duplex-link $LSR2 $LSR3 100Mb 10ms DropTail
$ns duplex-link $LSR3 $LSR4 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR5 100Mb 10ms DropTail
$ns duplex-link $LSR5 $node10 100Mb 10ms DropTail
$ns duplex-link $LSR1 $LSR6 100Mb 10ms DropTail
$ns duplex-link $LSR6 $LSR7 100Mb 10ms DropTail
$ns duplex-link $LSR7 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR8 $LSR9 100Mb 10ms DropTail
$ns duplex-link $LSR9 $LSR5 100Mb 10ms DropTail
$ns duplex-link $LSR2 $LSR7 100Mb 10ms DropTail
$ns duplex-link $LSR3 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR9 100Mb 10ms DropTail

#
# configure ldp agents on all mpls nodes
#
for {set i 1} {$i < 10} {incr i} {
    set a LSR$i
    for {set j [expr $i+1]} {$j < 10} {incr j} {
        set b LSR$j
```

```

        eval $ns LDP-peer $$a $$b
    }
    set m [eval $$a get-module "MPLS"]
    $m enable-reroute "new"
}

#configuram culori pentru diferitele pachete LDP
$ns ldp-request-color      blue
$ns ldp-mapping-color      red
$ns ldp-withdraw-color     magenta
$ns ldp-release-color      orange
$ns ldp-notification-color yellow

Classifier/Addr/MPLS enable-on-demand
Classifier/Addr/MPLS enable-ordered-control
[$LSR1 get-module "MPLS"] enable-data-driven

set src0 [new Application/Traffic/CBR]
set udp0 [new Agent/UDP]
$src0 attach-agent $udp0
$ns attach-agent $node0 $udp0
$src0 set packetSize_ 10000
$src0 set interval_ 0.001

set Dst0 [new Agent/LossMonitor]
$ns attach-agent $node10 $Dst0

$ns connect $udp0 $Dst0

$ns at 0.1 "$src0 start"

#protected LSP
$ns at 0.1 "[$LSR1 get-module MPLS] make-explicit-route 5 2_3_4_5 1000 -1"
$ns at 0.2 "[$LSR1 get-module MPLS] flow-erlsp-install 10 -1 1000"

#link failure
$ns rtmodel-at 0.3 down $LSR4 $LSR5

#restore link
$ns rtmodel-at 0.5 up $LSR4 $LSR5

$ns at 0.6 "$src0 stop"

$ns at 1.0 "finish"
$ns run

```


Fast_reroute.tcl

```
set ns [new Simulator]

set nf [open MPLSexample.nam w]
$ns namtrace-all $nf
set f0 [open MPLSexample.tr w]
$ns trace-all $f0

proc finish {} {
    global ns nf f0
    $ns flush-trace
    close $nf
    close $f0
    exec nam MPLSexample.nam &
    exit 0
}

#Agent/LDP set trace_ldp_ 1
#Classifier/Addr/MPLS set trace_mpls_ 1

$ns rtproto DV

set node0 [$ns node]
$ns node-config -MPLS ON
set LSR1 [$ns node]
set LSR2 [$ns node]
set LSR3 [$ns node]
set LSR4 [$ns node]
set LSR5 [$ns node]
set LSR6 [$ns node]
set LSR7 [$ns node]
set LSR8 [$ns node]
set LSR9 [$ns node]
$ns node-config -MPLS OFF
set node10 [$ns node]

$ns duplex-link $node0 $LSR1 100Mb 10ms DropTail
$ns duplex-link $LSR1 $LSR2 100Mb 10ms DropTail
$ns duplex-link $LSR2 $LSR3 100Mb 10ms DropTail
$ns duplex-link $LSR3 $LSR4 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR5 100Mb 10ms DropTail
$ns duplex-link $LSR5 $node10 100Mb 10ms DropTail
$ns duplex-link $LSR1 $LSR6 100Mb 10ms DropTail
$ns duplex-link $LSR6 $LSR7 100Mb 10ms DropTail
$ns duplex-link $LSR7 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR8 $LSR9 100Mb 10ms DropTail
$ns duplex-link $LSR9 $LSR5 100Mb 10ms DropTail
$ns duplex-link $LSR2 $LSR7 100Mb 10ms DropTail
$ns duplex-link $LSR3 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR9 100Mb 10ms DropTail

#
# configure ldp agents on all mpls nodes
#
for {set i 1} {$i < 10} {incr i} {
    set a LSR$i
    for {set j [expr $i+1]} {$j < 10} {incr j} {
        set b LSR$j
        eval $ns LDP-peer $$a $$b
    }
}
```

```

    }
    set m [eval $$a get-module "MPLS"]
    $m enable-reroute "new"
}

#configuram culori pentru diferitele pachete LDP
$ns ldp-request-color      blue
$ns ldp-mapping-color      red
$ns ldp-withdraw-color     magenta
$ns ldp-release-color      orange
$ns ldp-notification-color yellow

Classifier/Addr/MPLS enable-on-demand
Classifier/Addr/MPLS enable-ordered-control
[$LSR1 get-module "MPLS"] enable-data-driven
[$LSR3 get-module "MPLS"] enable-data-driven

set src0 [new Application/Traffic/CBR]
set udp0 [new Agent/UDP]
$src0 attach-agent $udp0
$ns attach-agent $node0 $udp0
$src0 set packetSize_ 10000
$src0 set interval_ 0.001

set Dst0 [new Agent/LossMonitor]
$ns attach-agent $node10 $Dst0

$ns connect $udp0 $Dst0

$ns at 0.1 "$src0 start"

#protected LSP
$ns at 0.1 "[$LSR1 get-module MPLS] make-explicit-route 5 2_3_4_5 1000 -1"
$ns at 0.1 "[$LSR1 get-module MPLS] flow-erlsp-install 10 -1 1000"

#LSR1's backup
$ns at 0.1 "[$LSR1 get-module MPLS] make-explicit-route 3 6_7_8_3 1001 -1"

#LSR2's backup
$ns at 0.1 "[$LSR2 get-module MPLS] make-explicit-route 4 7_8_4 1002 -1"

#LSR3's backup
$ns at 0.1 "[$LSR3 get-module MPLS] make-explicit-route 5 8_9_5 1003 -1"

#LSR4's backup
$ns at 0.1 "[$LSR4 get-module MPLS] make-explicit-route 5 9_5 1004 -1"

#link failure
$ns rtmodel-at 0.3 down $LSR4 $LSR5

#change flow to backup LSP
$ns at 0.3 "[$LSR4 get-module MPLS] flow-erlsp-install 10 -1 1004"

#restore link
$ns rtmodel-at 0.5 up $LSR4 $LSR5

$ns at 0.6 "$src0 stop"

$ns at 1.0 "finish"
$ns run

```

Propunere.tcl

```
set ns [new Simulator]

set nf [open MPLSexample.nam w]
$ns namtrace-all $nf
set f0 [open MPLSexample.tr w]
$ns trace-all $f0

proc finish {} {
    global ns nf f0
    $ns flush-trace
    close $nf
    close $f0
    exec nam MPLSexample.nam &
    exit 0
}

#Agent/LDP set trace_ldp_ 1
#Classifier/Addr/MPLS set trace_mpls_ 1

$ns rtproto DV

set node0 [$ns node]
$ns node-config -MPLS ON
set LSR1 [$ns node]
set LSR2 [$ns node]
set LSR3 [$ns node]
set LSR4 [$ns node]
set LSR5 [$ns node]
set LSR6 [$ns node]
set LSR7 [$ns node]
set LSR8 [$ns node]
set LSR9 [$ns node]
$ns node-config -MPLS OFF
set node10 [$ns node]

$ns duplex-link $node0 $LSR1 100Mb 10ms DropTail
$ns duplex-link $LSR1 $LSR2 100Mb 10ms DropTail
$ns duplex-link $LSR2 $LSR3 100Mb 10ms DropTail
$ns duplex-link $LSR3 $LSR4 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR5 100Mb 10ms DropTail
$ns duplex-link $LSR5 $node10 100Mb 10ms DropTail
$ns duplex-link $LSR1 $LSR6 100Mb 10ms DropTail
$ns duplex-link $LSR6 $LSR7 100Mb 10ms DropTail
$ns duplex-link $LSR7 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR8 $LSR9 100Mb 10ms DropTail
$ns duplex-link $LSR9 $LSR5 100Mb 10ms DropTail
$ns duplex-link $LSR2 $LSR7 100Mb 10ms DropTail
$ns duplex-link $LSR3 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR8 100Mb 10ms DropTail
$ns duplex-link $LSR4 $LSR9 100Mb 10ms DropTail

#
# configure ldp agents on all mpls nodes
#
for {set i 1} {$i < 10} {incr i} {
    set a LSR$i
    for {set j [expr $i+1]} {$j < 10} {incr j} {
        set b LSR$j
        eval $ns LDP-peer $$a $$b
    }
}
```

```

    }
    set m [eval $$a get-module "MPLS"]
    $m enable-reroute "new"
}

#configuram culori pentru diferitele pachete LDP
$ns ldp-request-color      blue
$ns ldp-mapping-color      red
$ns ldp-withdraw-color     magenta
$ns ldp-release-color      orange
$ns ldp-notification-color yellow

Classifier/Addr/MPLS enable-on-demand
Classifier/Addr/MPLS enable-ordered-control
[$LSR1 get-module "MPLS"] enable-data-driven

set src0 [new Application/Traffic/CBR]
set udp0 [new Agent/UDP]
$src0 attach-agent $udp0
$ns attach-agent $node0 $udp0
$src0 set packetSize_ 10000
$src0 set interval_ 0.001

set Dst0 [new Agent/LossMonitor]
$ns attach-agent $node10 $Dst0

$ns connect $udp0 $Dst0

$ns at 0.1 "$src0 start"

#protected LSP
$ns at 0.1 "[$LSR1 get-module MPLS] make-explicit-route 5 2_3_4_5 1000 -1"
$ns at 0.1 "[$LSR1 get-module MPLS] flow-erlsp-install 10 -1 1000"

#alternate LSP
$ns at 0.1 "[$LSR1 get-module MPLS] make-explicit-route 5 6_7_8_9_5 1001 -1"

#LSR2's backup
$ns at 0.1 "[$LSR2 get-module MPLS] make-explicit-route 7 7 1002 -1"

#LSR3's backup
$ns at 0.1 "[$LSR3 get-module MPLS] make-explicit-route 8 8 1003 -1"

#LSR4's backup
$ns at 0.1 "[$LSR4 get-module MPLS] make-explicit-route 9 9 1004 -1"

#link failure
$ns rtmodel-at 0.3 down $LSR4 $LSR5

#change flow to backup LSP
$ns at 0.3 "[$LSR4 get-module MPLS] flow-erlsp-install 10 -1 1004"

#restore link
$ns rtmodel-at 0.5 up $LSR4 $LSR5

$ns at 0.6 "$src0 stop"

$ns at 1.0 "finish"
$ns run

```