

**Políticas e Mecanismos de Engenharia de
Tráfego para Redes MPLS/DS**

Valéria Jotta dos Reis Paixão

Dissertação de Mestrado

Políticas e Mecanismos de Engenharia de Tráfego para Redes MPLS/DS

Valéria Jotta dos Reis Paixão

Fevereiro de 2007

Banca Examinadora:

- Prof. Dr. Edmundo R. M. Madeira (Orientador)
- Prof. Dr. Maurício Ferreira Magalhães (DCA/FEEC/Unicamp)
- Prof. Dr. Nelson Luis Saldanha da Fonseca (IC/Unicamp)
- Prof. Dr. Luiz Eduardo Buzato - Suplente (IC/Unicamp)

**FICHA CATALOGRÁFICA ELABORADA PELA
BIBLIOTECA DO IMECC DA UNICAMP**
Bibliotecária: Maria Júlia Milani Rodrigues – CRB8a / 2116

Paixão, Valéria Jotta dos Reis

P167p Políticas e mecanismos de engenharia de tráfego para redes
MPLS/DS / Valéria Jotta dos Reis Paixão - Campinas, [S.P. :s.n.], 2007.

Orientador : Edmundo Roberto Mauro Madeira

Trabalho final (mestrado profissional) - Universidade Estadual de
Campinas, Instituto de Computação.

1. Internet (Redes de computação). 2. Programação
(Computadores) – Gerência. 3. Engenharia de tráfego. I. Madeira,
Edmundo Roberto Mauro. II. Universidade Estadual de Campinas.
Instituto de Computação. III. Título.

Título em inglês: Policies and traffic engineering mechanisms for MPLS/DS networks

Palavras-chave em inglês (Keywords): 1. Internet (Computer network). 2. Programming
(Computer programming) – Management. 3. Traffic engineering.

Área de concentração: Redes de Computadores

Titulação: Mestre em Ciência da Computação

Banca examinadora: Prof. Dr. Edmundo Roberto Mauro Madeira (IC/UNICAMP)
Prof. Dr. Maurício Ferreira Magalhães (DCA/FECC/UNICAMP)
Prof. Dr. Nelson Luis Saldanha da Fonseca (IC/UNICAMP)
Prof. Dr. Luiz Eduardo Buzato (IC/UNICAMP)

Data da defesa: 23/02/2007

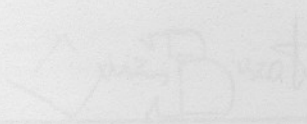
Programa de Pós-Graduação: Mestrado em Ciência da Computação

“Políticas e Mecanismos de Engenharia de Tráfego para Redes MPLS/DS”

Trabalho Final Escrito defendido e aprovado em 23 de fevereiro de 2007, pela
Banca Examinadora composta pelos Professores Doutores:

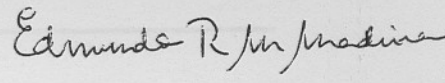
Este exemplar corresponde à redação final do
Trabalho Final devidamente corrigida e defendida
por Valéria Jotta dos Reis Paixão e aprovada pela
Banca Examinadora.


Prof. Dr. Mauricio Ferreira Magalhães
FEEC - UNICAMP


Prof. Dr. Luiz Eduardo Suzato
IC - UNICAMP

Campinas, 23 de fevereiro de 2007.


Prof. Dr. Edmundo Roberto Mauro Madeira
IC - UNICAMP

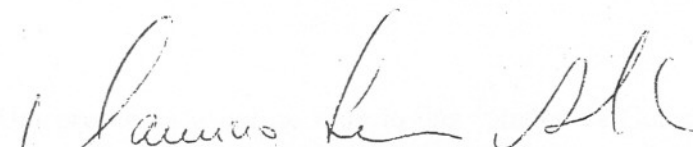

Prof. Dr. Edmundo Roberto Mauro Madeira
(Orientador)

Trabalho Final apresentado ao Instituto de
Computação, UNICAMP, como requisito
parcial para a obtenção do título de
Mestre em Computação na área
de Engenharia de Computação

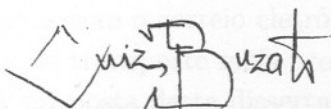
2008 06364

TERMO DE APROVAÇÃO

Trabalho Final Escrito defendido e aprovado em 23 de fevereiro de 2007, pela
Banca Examinadora composta pelos Professores Doutores:



Prof. Dr. Mauricio Ferreira Magalhães
FEEC - UNICAMP



Prof. Dr. Luiz Eduardo Buzato
IC - UNICAMP



Prof. Dr. Edmundo Roberto Mauro Madeira
IC - UNICAMP

Resumo

Um problema atual no mundo das “Redes de Computadores” é o fornecimento de QoS fim-a-fim na Internet a diferentes usuários, ou seja, a distintos tipos de serviços. Além disso, a demanda por este serviço só vem aumentando uma vez que a Internet está sendo cada vez mais solicitada não somente para transferência de arquivos, *web* ou para tarefas relacionadas com o correio eletrônico, mas também para aplicações multimídia que solicitam meios de transporte mais precisos.

A proposta desta dissertação é fornecer melhorias a uma plataforma MPLS-DiffServ, transformando a mesma em um sistema mais flexível e dinâmico no oferecimento de qualidade de serviço fim-a-fim. As novas características utilizam tipos de tráfegos amplamente encontrados na Internet, assim como dois parâmetros de monitoramento, a largura de banda do LSP e os níveis de descartes de pacotes.

As inovações vêm no formato de novas políticas para o gerenciamento da rede em caso da ruptura dos acordos pré-estabelecidos com os clientes, além de uma nova estrutura introduzida no mecanismo de Engenharia de Tráfego que ajudará o sistema a fazer uma análise mais profunda do tráfego no LSP de modo que as decisões mais exatas e rápidas possam ser aplicadas a fim de recuperar circunstâncias ideais para atender às expectativas dos usuários finais.

Na mesma plataforma foi feita a verificação das melhorias para as quais o sistema forneceu respostas bem sucedidas que estimulasse a ambição de fazer da plataforma uma opção de ferramenta para companhias que desejem simular novas políticas administrativas antes de executá-las em uma rede real que afetam usuários reais.

Abstract

A current problem in the “Computer Network“ world is to provide end-to-end QoS in the Internet to different users what means supply support to distinct types of services. Moreover, the demand for this support is increasing even more once the Internet is being handled not only for transferring files, accessing the web or for email tasks but also for multimedia applications, it means, services that require a more accurate transmission media.

The proposal of this dissertation is to provide enhancements to an MPLS-DiffServ platform, making it up a more flexible and dynamic system to offer end-to-end QoS. The new features make use of types of traffic largely handled in the Internet as well as two monitor parameters, the LSP bandwidth and the levels of packet drop.

The innovation comes in the format of new policies in order to manage the networks in case of break of agreements made with the clients, besides a new structure that was inserted in the “Traffic Engineering” mechanism that will help in a sense of deep analysis of the traffic in the LSP so that more accurate and fast decisions can be applied in order to recover ideal conditions to attend the expectations of the final users.

On this framework the verification of the new improvements was done and, as a result, the system provided a successful response that copes with the ambition of making the platform an option tool for companies that wish to simulate new administrative policies just before to implement them in a real network that affects real users.

Agradecimentos

Primeiramente, eu gostaria de agradecer a minha grande amiga Nossa Senhora pela companhia inseparável durante essa jornada. Ela que sempre intercedeu por mim junto a Papai do Céu para que colocasse em meu caminho pessoas especiais em mais esse passo da minha vida.

Dentre essas pessoas que muito fazem diferença na minha existência, eu devo mencionar os meus pais, que apesar de humildes, sempre demonstraram uma grande consciência diante da importância do conhecimento para a formação de um indivíduo, o que me guiou até essas palavras. Eu poderia dedicar um capítulo inteiro sobre a minha mãe e sobre o meu pai, mas não conseguiria expressar em palavras a grande força e inspiração que eles representam na minha vida.

Em seguida, eu gostaria de mencionar o meu marido e meu filho que me proporcionam um lar muito feliz e, através dessa felicidade, surge empolgação e tranquilidade para realizar esse trabalho.

Com relação a respeito e compreensão dispensada durante essa etapa da minha vida, eu posso falar do meu orientador Prof. Dr. Edmundo Madeira que muito me ajudou e muito me ensinou com seu profissionalismo.

Por último, porém não menos agradecida, deixo aqui registrado o meu agradecimento aos meus colegas de laboratório Patrick e Leonel que me ajudaram de uma maneira memorável nessa fase final da dissertação, com dicas e bate-papos valiosíssimos, assim como a Cláudia da secretaria que sempre se mostrou muito prestativa.

"Nós nunca estamos só, Deus sempre está conosco."

(Pedrina Jotta - minha adorada mãe)

"O saber é um tesouro que ninguém consegue tirar de uma pessoa."

(João da Volta - meu bisavô)

"Vás por onde conheces, mesmo que este seja o caminho mais longo."

(Paulo Américo da Silveira Jotta - meu amado avô)

"Tudo vale a pena se a alma não é pequena."

(Fernando Pessoa)

"Quem sabe faz a hora, não espera acontecer."

(Geraldo Vandré)

Sumário

Resumo	v
Abstract	vii
Agradecimentos	ix
Acrônimos	xvii
1 Introdução	1
2 Fundamentos Teóricos	5
2.1 Qualidade de Serviço	5
2.1.1 Serviços Diferenciados (DiffServ)	6
2.1.2 Identificação e Classificação de Pacotes	7
2.1.3 Policiamento e condicionamento de tráfego	7
2.1.4 DS CodePoint - DSCP	8
2.1.5 Comportamento por Hop	9
2.2 Troca de Rótulos Multiprotocolar	13
2.2.1 Classe Equivalente de Encaminhamento	13
2.2.2 Rótulos	14
2.2.3 Agregação de Tráfego	15
2.2.4 Seleção de Rota	15
2.3 Engenharia de Tráfego	15
2.4 MPLS e a Engenharia de Tráfego	17
2.5 MPLS e Serviços Diferenciados	18
2.6 Controle de admissão	19
2.7 Políticas - gerenciamento de sistemas	19
2.7.1 Tipos de Políticas	19
2.7.2 Gerenciamento de Políticas	20

3	Trabalhos Relacionados	23
3.1	Gerenciamento de Tráfego	23
3.2	Plataformas	25
4	Plataforma Proposta	27
4.1	Network Simulator	27
4.1.1	MPLS-DS	29
4.1.2	Integração DS-NORTEL & MNS	34
4.1.3	Agente TCP/SinkMonitor	34
4.2	Arquitetura do Bandwidth Broker	35
4.2.1	Gerente de Recursos	36
4.2.2	Gerente de Políticas	37
4.2.3	Gerente de negociação	37
4.2.4	Objeto de Configuração	37
4.2.5	Monitoramento de Tráfego Simulado	37
4.2.6	Objeto de Aplicação	38
4.2.7	Objeto de Gerenciamento	38
4.2.8	Repositório de Dados e Contratos de Serviços	38
4.3	Arquitetura de Políticas	40
4.3.1	Políticas de Inicialização	40
4.3.2	Políticas de Atuação	41
4.3.3	Políticas Pré-existentes	41
4.3.4	Reformulação/Criação de Políticas de Atuação	42
4.3.5	Políticas de Negociação	45
5	Experimentos/Validações	47
5.1	Experimento 1	48
5.2	Experimento 2	51
5.3	Experimento 3	55
5.4	Considerações Finais sobre as Simulações Realizadas	59
6	Conclusão e Trabalhos Futuros	61
6.1	Conclusão	61
6.2	Trabalhos Futuros	62
	Referências	64

Lista de Tabelas

2.1	<i>Codepoints</i> para Descarte e Classe de Encaminhamento do Grupo AF . . .	10
5.1	Características dos tráfegos gerados	48
5.2	Características das classes de serviço	48
5.3	Tabela LFT	49
5.4	Tabela LFT	49
5.5	Tabela LFT - tráfego FTP	52
5.6	Tabela LFT - tráfegos FTP e Voz	52
5.7	Tabela LFT - Adição do primeiro fluxo da simulação (experimento 3) . . .	55
5.8	Tabela LFT - Adição do segundo fluxo da simulação (experimento3)	56

Lista de Figuras

2.1	Arquitetura de Política do IETF/DMTF	20
4.1	Visão Simplificada do NS	28
4.2	Estrutura de Tabelas do MNS	30
4.3	Classe MPLDSAddressClassifier	32
4.4	Estrutura de Tabelas do MNS incluindo a Tabela LFT	34
4.5	Arquitetura do Bandwidth Broker	36
5.1	Topologia de rede usada nos experimentos	47
5.2	Porcentagem de ocupação de banda para os LSPs 2000 e 2999	50
5.3	Medidas de atraso dos cliente 01 ao Cliente 03 e do cliente 02 ao cliente 04	51
5.4	Vazão do tráfego FTP	53
5.5	Número de Pacotes Recebidos, Duplicados e Perdidos	53
5.6	Descarte de pacotes para o LSP2000 - classe de serviço EF	54
5.7	Níveis de ocupação de banda para os LSPs 2000 e 2999	54
5.8	Vazão do tráfego FTP	56
5.9	Vazão do tráfego FTP	57
5.10	Descarte de pacotes para o LSP2000 - classe de serviço EF	57
5.11	Atraso de pacotes para o LSP2000 - classe de serviço EF	58
5.12	Níveis de ocupação de banda para os LSPs 2000 e 2999	58

Acrônimos

AF	Assured Forwarding
ATM	Asynchronous Transfer Mode
BGRPP	Border Gateway Reservation Protocol Plus
BRR	Bit-by-bit Round Robin
CBQ	Class-Based Queuing
CBR	Constant Bit Rate
COPS	Common Open Policy Service
CORBA	Common Object Request Broker Architecture
DiffServ	DIFferential SERVices
DII	Dynamic Invocation Interface
DSI	Dynamic Skeleton Interface
DS-TE	Differentiated Service - Traffic Enginnering
DWDM	Dense Wavelength Division Multiplexing
E-LSP	EXP-inferred-PSC LSPs
ERB	Explicit Routing Information Base

FEC	Forwarding Equivalence Class
FIFO	First-In-First-Out
IP	Internet Protocol
LDAP	Lightweight Directory Access Procol
LDP	Label Distribution Protocol
LIB	Label Information Base
LSP	Label Switched Path
L-LSP	Label-Only-Inferred-PSC
LSR	Label Switch Router
MPLS	MultiProtocol Label Switching
OSFP	Open Shortest Path First
OTCL	Object-Oriented TCL
PDP	Policy Decision Point
PEP	Policy Enforcement Point
PFT	Partial Forwarding Table
PHB	Per-Hop-Behavior
PQ	Priority Queue
QoS	Quality Of Service
RED	Random Early Detection

SIBBS	Simple Interdomain Bandwidth Broker Signalling
SNMP	Simple Network Management Protocol
TCL	Tool Control Language
TCP	Transmission Control Protocol
TOS	Type of Service
UDP	User Datagram Protocol
VoIP	Voice Over Internet Protocol
VBR	Variable Bit Rate
WFQ	Weighted Fair Queue
WRR	Weighted Round Robin

Capítulo 1

Introdução

Encontra-se a motivação e a importância desse trabalho na melhoria do desempenho da Internet no atendimento das necessidades das aplicações multimídia.

A Internet é uma rede de multi-serviços que provê infra-estrutura de comunicação universal para negócios e lazer, onde a qualidade de serviço (QoS - *Quality Of Service*) e a Engenharia de Tráfego [AMA⁺99] tornaram-se atributos essenciais devido aos novos requisitos de seus usuários e aplicações.

Com o crescimento vertiginoso do número de usuários da Internet, precisa-se repensar na eficiência das funcionalidades da família TCP/IP (*Transmission Control Protocol/Internet Protocol*). A arquitetura TCP/IP foi concebida para oferecer um serviço de rede do tipo melhor-esforço. Nesse sentido, os protocolos de rede IP e os de transporte TCP e UDP (*User Datagram Protocol*) fornecem mecanismos muito limitados de qualidade de serviço que seriam o decréscimo multiplicativo e a partida lenta. Porém, não se pode deixar de mencionar, que os mecanismos citados atendem às aplicações tradicionais, como acesso *web*, correio eletrônico e transferência de arquivos, pois são serviços que são passíveis a atrasos.

Vários estudos têm se destinado na melhoria da QoS para a Internet no sentido de atender a demanda de aplicações cujo tráfego seja de alta taxa de transmissão com baixa tolerância a atrasos e também precisem de um gerenciamento de segurança, como é o caso de aplicações de áudio e vídeo que trafeguem na Internet e/ou em ambientes comerciais. O modelo de Troca de Rótulos Multiprotocolar (MPLS - *MultiProtocol Label Switching*) integrado ao modelo de Serviços Diferenciados (Diffserv- *DIFFerential SERVices*) [FWD⁺03], complementado pelo mecanismo de políticas de comportamentos e admissão de tráfego, se apresenta como uma solução interessante para a implementação de aplicações.

Por um lado, tem-se o MPLS que se apresenta como uma solução versátil para a solução de problemas relacionados à velocidade de encaminhamento, Engenharia de Tráfego, gerenciamento eficiente de largura de banda, podendo ser implantada sobre ATM (*Asynch-*

ronous Transfer Mode), DWDM (*Dense Wavelength Division Multiplexing*), Frame Relay, Ethernet etc.

A Engenharia de Tráfego facilita e torna confiáveis as operações de rede enquanto, simultaneamente, otimiza a utilização de recursos da rede e o desempenho do tráfego.

A utilização da tecnologia MPLS [RVC01] possibilita a sua extensão para controle das futuras redes óticas de transporte (*Multiprotocol Lambda Switching*) e coexistência com as redes IP atuais, possibilitando uma implantação com incremento.

Do outro lado, o DiffServ é um mecanismo elegante que provê uma maneira simples de fornecer diferentes classes de serviços para os dados que trafegam na Internet mediante uma prévia negociação, podendo ser implantado em larga escala.

A estrutura complementar de políticas acompanha a plataforma MPLS-DiffServ no sentido de fornecer, em um primeiro passo, um refinamento no atendimento aos requisitos de QoS para os clientes da rede, e, em um segundo passo, a propriedade de adaptação dinâmica de comportamento sem a necessidade de parada ou gravação das propriedades do sistema.

Este trabalho faz uso de uma plataforma MPLS-DS [Gui01] que incorpora ao *backbone* Internet as habilidades de diferenciar serviços e aplicar engenharia de tráfego por meio do protocolo MPLS. Os recursos são disponibilizados aos clientes através de um modelo dinâmico e automático de negociação de serviços por meio de contratos, podendo se dar não apenas no interior do domínio, mas também entre domínios distintos.

Além disso, a plataforma possui um sistema de gerenciamento de recursos de rede baseado em um *Bandwidth Broker* [Arc01] que interage com uma arquitetura de políticas, visando manter os contratos pré estabelecidos pelos clientes que fazem uso do backbone, sendo este o campo de atuação desta trabalho.

O objetivo desta dissertação é oferecer um conjunto confiável de políticas para aplicações envolvendo tráfego de áudio e vídeo tais como videoconferência, tele-educação e vídeo por demanda, bem como tráfego de Voz sobre IP (VoIp - *Voice Over Internet Protocol*). Além de tornar disponível um novo mecanismo ao controle de tráfego, aumentando os atributos de “Engenharia de Tráfego” já existentes na plataforma original.

As políticas implementadas atuam dinamicamente após a assinatura do contrato, assegurando requisitos de descarte de pacotes e alocação de recursos, mesmo estando suscetível a alterações de disponibilidade de recursos.

Esta dissertação também irá levar em consideração, o momento da admissão do tráfego (políticas de admissão) para que sejam cumpridos pré requisitos que garantam a eficiência das políticas disponíveis.

A dissertação encontra-se estruturada da seguinte forma: o Capítulo 2 discorre a parte teórica introdutória englobando os temas de QoS, DiffServ, Engenharia de Tráfego, assim como Arquitetura de Políticas e demais temas relacionados, o Capítulo 3 apresenta

trabalhos relacionados com as políticas de gerenciamento de tráfego e as plataformas que foram utilizadas para validar este dados, enquanto que o Capítulo 4 tem a finalidade de expor a plataforma utilizada na verificação e validação das políticas propostas, bem como explica a nova estrutura que irá auxiliar na “Engenharia de Tráfego” da mesma. O Capítulo 5 traz os dados práticos obtidos através da simulações feitas sobre a plataforma mencionada no Capítulo 4. Finalmente, a conclusão da dissertação é exposta no Capítulo 6 onde também se encontram as propostas de trabalhos futuros.

Capítulo 2

Fundamentos Teóricos

Neste capítulo serão abordados os aspectos teóricos envolvidos na dissertação e que estão relacionados à integração dos conceitos de Qos, DiffServ, Engenharia de Tráfego e Arquitetura de Políticas.

2.1 Qualidade de Serviço

A Internet é capaz de oferecer um serviço caracterizado como melhor-esforço. O melhor-esforço é um mecanismo que visa realizar o máximo de esforço para prestar o serviço para o qual foi projetado, ou seja, transmissão de dados da origem ao destino, estando os dados sujeitos a flutuações de velocidade, banda e atraso.

O conceito de “Qualidade de Serviço” tem como objetivo a inserção de novas possibilidades de respostas para os diferentes tipos de tráfego que podem acessar a rede, sendo este o seu grande desafio, ou seja, prover uma arquitetura que seja capaz de sustentar uma resposta desejada de serviço de ponta a ponta diante da grande diversidade de plataformas adotadas na Internet.

Para oferecer essas distinções entre os serviços existe não só uma requisição para prover resposta a serviços diferenciados, como também uma requisição para controlar a carga de serviço qualificado na rede, de maneira que os recursos alocados para garantir uma resposta a um serviço em particular são capazes de prover a resposta para a carga imposta.

Além disso, o termo “Qualidade de Serviço” está intimamente associado às redes orientadas a conexão. De fato, cada equipamento por ocasião do estabelecimento da conexão reserva recursos para o processamento de pacotes que fluem pela conexão, sendo estes recursos tipicamente associados ao armazenamento e processamento de pacotes.

Enquanto que as redes não orientadas a conexão, como as redes IP, processam os pacotes de forma individual, não armazenando recursos para o tráfego que necessita de

garantias de Qualidade de Serviço. Portanto, prover QoS em redes não orientadas à conexão exige a definição de fluxos para que recursos sejam reservados com a finalidade de oferecer um tratamento adequado aos pacotes que fluem através das mesmas. Os fluxos são obtidos através da agregação de pacotes com alguma relação entre si. Esta agregação pode se dar em função da definição de rotas como no MPLS.

O número de fluxos em roteadores do tipo núcleo pode inviabilizar a introdução de qualidade de serviço na Internet pela incapacidade dos roteadores:

- armazenarem informações sobre o estado de cada fluxo;
- computarem o tratamento a ser conferido aos pacotes associados a cada fluxo;
- atualizarem o estado de cada fluxo em função do seu estabelecimento, encerramento, bem como em ações de roteamento;

Para resolver estas questões apresentadas acima uma forma mais simples de prover qualidade de serviço foi proposta. A idéia é fazer uma agregação de diversos fluxos em um número reduzido de classes de serviços (CoS - *Class of Service*).

A agregação de milhares de milhares de fluxos em algumas poucas classes impede a definição individual de parâmetros de QoS para cada fluxo. Como roteadores não mantêm os estados de fluxos, esta estratégia exige que cada pacote carregue consigo a classe a qual pertence.

As Classes de Serviço podem ser implementadas de duas maneiras:

- Atribuindo-se prioridade às classes: cada classe possui uma prioridade associada, relativa às demais classes. Pacotes de prioridades mais altas possuem tratamento privilegiado em relação aos de prioridade mais baixa.
- Através de comportamento padrão por *hop* (*per-hop behavior* - PHB): esta estratégia define um comportamento padrão dispensado à classe por parte dos roteadores. PHB resgata, de certa forma, o campo TOS do protocolo IPv4. A arquitetura de serviços diferenciados emprega esta estratégia que é a abordada neste trabalho.

2.1.1 Serviços Diferenciados (DiffServ)

No contexto de transmissão de pacotes, entende-se por "Serviço" algumas características da transmissão de um pacote através de uma rede. Essas características podem ser dimensionadas por meio do atraso, *jitter*, descarte ou mesmo prioridade de acesso aos recursos da rede. Uma diferenciação nos serviços se faz necessária para atender aos requisitos

determinísticos de alguma aplicação ou usuário. Dessa maneira, faz-se possível a implementação de QoS por meio do tratamento específico do tráfego durante o seu trajeto na rede.

A arquitetura definida para a implantação de Serviços Diferenciados tem seu foco no encaminhamento baseado no comportamento por hop (PHB - per hop forwarding behavior) [NC01]. Porém, sem deixar de levar em consideração funções de classificação de pacotes, assim como funções de condicionamento de tráfego, incluindo medição, marcação, formatação e policiamento.

2.1.2 Identificação e Classificação de Pacotes

A Identificação de pacotes nos equipamentos de borda visa associar o pacote a um fluxo.

Cinco parâmetros são comumente empregados para associar pacotes à fluxos:

- Endereço IP de origem
- Endereço IP de destino
- Protocolo de camada superior (TCP, UDP, etc)
- Porta TCP ou UDP de origem
- Porta TCP ou UDP de destino

Recebido um datagrama, o equipamento de borda deve associá-lo a um fluxo, tendo como base os 5 parâmetros acima ou determinar que o pacote não está associado a nenhum fluxo. Após a identificação do fluxo e roteamento o pacote é classificado segundo a política de escalonamento empregada. Via de regra a classificação se resume em posicionar o pacote em determinada fila da interface de saída.

2.1.3 Policiamento e condicionamento de tráfego

Policiamento de tráfego visa garantir que fontes de tráfego não violem parâmetros de QoS estabelecidos para seus fluxos, visando homogeneizar os fluxos que ingressam na rede evitando que surtos de tráfego causem congestionamento nos equipamentos núcleos. Policiamento e condicionamento de tráfego são atividades processadas de forma integrada.

O algoritmo *token-bucket* é um algoritmo empregado para as funções de policiamento e condicionamento de tráfego [Tan04] usado para várias aplicações, onde é melhor permitir uma saída rápida mesmo na presença de rajadas de pacotes.

O algoritmo "token bucket" nos permite configurar o número de pacotes no balde e permite que rajadas com um tamanho até "n" possam ser enviadas. O comprimento

do *token-bucket* é ajustado de acordo com o descritor de tráfego de fonte e de recursos disponíveis no equipamento, principalmente *buffers* disponíveis nas interfaces de rede. O descritor da fonte é composto da taxa de pico (bits/s), taxa sustentada (bits/s) e comprimento da rajada (bits). O algoritmo *token-bucket* opera do seguinte modo para cada fluxo sendo policiado. Fichas (*tokens*) são depositadas num balde com capacidade C para fichas a taxas constantes. A transmissão de um pacote armazenado em sua interface de saída consome uma ficha do balde. Caso o balde esteja vazio, o pacote no topo a fila é considerado não conforme com os parâmetros de tráfego especificado. Neste caso, o envio do pacote é atrasado, matendo-o na fila até que uma ficha seja depositada no balde; ou simplesmente, o mesmo é descartado.

Numa situação de regime, o número de fichas no balde permanece constante, pois tanto pacotes quanto fichas chegam a uma taxa constante. Caso a fonte pare a transmissão por um período, as fichas vão se acumulando no balde. Numa situação de rajadas espaçadas, a fila de saída poderá ter a capacidade excedida. Neste caso, o descarte de pacotes é a única alternativa.

O algoritmo *token-bucket* ainda efetua condicionamento de tráfego pois impede que rajadas superiores a um determinado comprimento sejam transmitidas. Entretanto, a função de policiamento de tráfego apresenta um problema.

O problema se refere ao caso em que o balde está totalmente cheio e uma rajada ocorre. Neste cenário, como existem fichas em quantidade, toda a rajada pode ser transmitida, porém esta transmissão pode ocupar toda a capacidade da banda disponível, impedindo que pacotes pertencentes a outros fluxos sejam transmitidos.

O algoritmo *leaky-bucket* se apresenta como uma solução para a questão abordada acima, pois ao ser colocado na saída do produto do algoritmo *token-bucket*, o mesmo se comporta como deste segundo balde de onde os pacotes são retirados a uma taxa não superior à taxa de pico estabelecida para o fluxo. Assim, a rajada oriunda do algoritmo *token-bucket* é transmitida a uma taxa inferior à capacidade do enlace.

Em outras palavras, se um ou mais processos no *host* tentam enviar um pacote quando o número máximo já foi atingido, o pacote é automaticamente descartado. Este acordo pode ser construído na interface de "hardware" ou simulado pelo sistema operacional do *host*. Na verdade, o algoritmo *leaky-bucket* funciona como um sistema servidor de enfileiramento com a tempo de serviço constante.

2.1.4 DS CodePoint - DSCP

O *DSCP*[BCB06] é um número que pode variar de 0 até 63. Este número é colocado no pacote IP para marcar a classe de tráfego a qual o mesmo pertence. Metade desses valores são valores padrões e a outra metade está disponível para marcações locais.

2.1.5 Comportamento por Hop

“Comportamento por Hop” (PHBs - *Per-Hop-Behavior*) define o comportamento que um tráfego irá seguir durante o encaminhamento na rede. Essa definição de comportamento pode se dar por meio da prioridade de recursos perante os outros PHBs ou em termos das características do tráfego, como atraso e descarte.

Esse mecanismo de identificação de um conjunto de tráfego se faz importante na presença de competição pelos recursos de um nó, pois estabelece os critérios de prioridade de acesso aos recursos.

O PHB é selecionado no nó pelo mapeamento do DSCP na recepção do pacote, sendo o DS um campo do cabeçalho IP e o *codepoint* os seus seis primeiros bits sendo que PHBs padrões possuem um *codepoint* recomendado. O mapeamento entre PHBs e *codepoints* pode ser um-para-um ou muitos-para-um. No primeiro caso um PHB está associado a um único *codepoint* enquanto no segundo caso um PHB está associado a vários *codepoints*.

Ligado aos PHBs estão os grupos de PHB que são implementados simultaneamente, levando em consideração limitações comuns relacionadas aos PHBs, tais como os serviços de fila, assim como a política de gerenciamento das mesmas. Um único PHB é um caso especial de um grupo de PHB.

Existem definidos quatro tipos de grupos de PHBs: PHB Padrão, Encaminhamento Garantido, Encaminhamento Expresso e o Seletor de Classes.

PHB Padrão

Na ausência de um mapeamento conhecido é atribuído ao pacote um PHB padrão, que não receberá tratamento diferenciado ou preferencial. Tipicamente, o PHB padrão é destinado ao serviço de melhor-esforço. O valor recomendado de *codepoint* para esta categoria de PHB é '000000'.

PHB Encaminhamento Garantido

Este tipo de grupo de padrão de PHB atende às necessidades de um usuário ou aplicação que necessita de uma alta garantia de encaminhamento dos pacotes IP na Internet.

O Encaminhamento Garantido é um tipo de comportamento de encaminhamento com algum nível de recursos de fila associados e três precedências de descartes [Gro02]. O grupo AF define quatro classes de encaminhamento. No grupo AF (*Assured Forwarding*) o serviço provido por uma classe é independente das demais classes, sendo função apenas dos recursos alocados pelo domínio para a classe e do tráfego agregado na classe.

Em cada classe AF os pacotes IP são marcados com uma das três possibilidades de precedência de descarte. Em caso de congestionamento, a precedência de descarte do pacote determina qual a importância do pacote dentro do contexto da classe AF. Um nó

congestionado descarta, primeiramente, os pacotes com um valor maior de precedência de descarte, no caso de um longo período em congestionamento. No caso de um período curto de congestionamento, os pacotes são enfileirados.

A Tabela 2.1 apresenta uma sugestão de *codepoints* associados a cada classe de encaminhamento e precedência de descarte para o grupo AF [HBWW99].

Descarte	Classe1	Classe2	Classe3	Classe4
Baixa	001010	010010	011010	100010
Média	001100	010100	011100	100100
Alta	001100	010100	011110	100110

Tabela 2.1: *Codepoints* para Descarte e Classe de Encaminhamento do Grupo AF

Encaminhamento Expresso

O grupo PHB denominado Encaminhamento Expresso (*Express Forward* - EF) [DCB⁺02] é uma solução para os serviços que desejem uma baixa taxa de perda, latência e *jitter*. Os pacotes que são marcados com esse PHB devem encontrar suas filas de encaminhamento, de preferência, vazias ou quase vazias, para isso é necessário garantir que a taxa de serviço de pacotes exceda a taxa de chegada dos mesmos em intervalos longos ou curtos, independentemente, da carga dos outros tráfegos não EF.

A configuração dos DS-capazes para o provimento deste PHB devem satisfazer a dois requisitos:

- O tráfego agregado neste PHB deve ter uma taxa de partida mínima (atraso interno) independente de outros fluxos que trafeguem pelo nó.
- O tráfego agregado neste PHB deve ser condicionado de forma que a taxa de ingresso em qualquer nó seja sempre inferior à taxa mínima de egresso do nó. Esta configuração é necessária para manter as filas nas interfaces de saída sempre na situação próxima de vazio.

Outra propriedade deste PHB é evitar o descarte para o tráfego em conformidade. Como as filas de saída não sofrem exaustão, descarte por *overflow* de *buffer* não deve ocorrer. Ainda como forma de evitar o descarte, este PHB proíbe a remarcação de pacotes no interior do domínio, ou seja, um datagrama que entra neste domínio com este PHB circula por todo o domínio com este PHB. PHB EF também é severo para o tráfego não conforme: simplesmente o descarta.

O PHB Encaminhamento Expresso tem o *codepoint* igual a 101110.

Seletor de Classe

Para manter a compatibilidade entre as redes DiffServ, campo DSField, e as redes que utilizam o campo Precedência do *byte* TOS (*Type of Service*) o cabeçalho IP, que tem como função marcar a prioridade do tráfego, criou-se o grupo PHB seletor de classe.

Os *codepoints* para o Grupo Seletor de Classe assumem o formato 'xxx000', onde os três primeiros bits são os bits de precedência do cabeçalho IPv4. Dessa forma, se um pacote é recebido de uma rede não DiffServ que usa o campo de precedência do TOS, a roteador vai entender como um PHB Seletor de Classe que, caso seja diferente do PHB padrão, deve receber um tratamento adequado e preferencial.

Associada aos diferentes tipos de grupos de PHB temos as políticas de escalonamento de filas que se baseiam em prioridade e as políticas de congestionamento para dar suporte à implementação de tais configurações.

Políticas de Escalonamento de Filas

Em casos de escassez de recursos para atender a uma demanda significativa é necessário o uso de filas de espera. As políticas de gerenciamento dessas filas atendem pelo nome de “políticas de escalonamento” que se responsabilizam por:

- tempo de espera finito;
- otimização de tempo de resposta, vazão, etc;
- garantida de progressão do sistema;
- evitar ociosidade e indisponibilidade crítica de recursos.

Como políticas de escalonamento podemos citar: FIFO (*First-In-First-Out*), PQ (*Priority Queue*), WRR (*Weighted Round Robin*) e WFQ (*Weighted Fair Queue*).

- FIFO

Esta política de escalonamento é uma política de implementação simples, onde os datagramas de um mesmo fluxo são transmitidos na mesma ordem de chegada pela fila de saída. Em caso de não haver espaço para empilhamento de um pacote na fila de saída, o mesmo é, automaticamente, descartado.

Um ponto negativo dessa política é a sua incapacidade de diferenciar fluxos, ou seja, pacotes com alta prioridade são tratados da mesma maneira que pacotes de baixa prioridade. Dessa forma, esta política de escalonamento não dá suporte a Serviços Diferenciados.

- PQ

A política de escalonamento é baseada em prioridades e suporta Serviços Diferenciados, pois ao contrário da FIFO possui várias filas na interface de saída. Essas filas da interface de saída possuem cada qual a sua prioridade, que são servidas em ordem decrescente de prioridade, ou seja, as filas de maior prioridade são atendidas, em seguida, das filas de menor prioridade.

Dessa maneira, apresenta as seguintes desvantagens:

1. pacotes nas filas de menor prioridade podem nunca ser atendidos, ou mesmo, ter um tempo de espera exacerbado;
2. métricas de medidas de atrasos podem ser facilmente violadas;
3. as prioridades associadas às filas não são adaptáveis ao comportamento da rede;
4. de forma individual, as filas PQ se comportam como filas FIFO.

- WRR

Esta política de escalonamento estabelece um percentual de ocupação da banda de acordo com a prioridade, ou seja, as filas continuam possuindo prioridade e continuam sendo atendidas na ordem decrescente, como a política PQ, porém a fila de menor prioridade irá esperar para ser atendida por um tempo definido, de acordo com o peso da fila de maior prioridade, e não até a fila de maior prioridade ser esvaziada. Dessa maneira, problemas relacionados com espera infinita, ou mesmo, por um tempo demasiadamente alto são evitados.

Com relação aos demais problemas relacionados para a política PQ, pode-se afirmar que eles continuam a existir.

- WFQ

WFQ é uma política de escalonamento que implementa o algoritmo BRR (*Bit-by-bit Round Robin*) sem a restrição de servir as filas bit-a-bit.

O BRR é um algoritmo de escalonamento que estipula que cada fluxo seja mantido em uma fila de saída dedicada e a cada ciclo do escanador, um bit de cada fila é enviado pela rede, o que o torna inviável. Dessa maneira, a proposta do WFQ é a sua implementação de acordo com a seguinte adaptação: ao receber um pacote, calcula-se o tempo de envio do último bit do pacote, caso o BRR estivesse sendo implementado na íntegra. Após esse cálculo o pacote é colocado na única fila de saída que está ordenada por tempo de partida do último bit. Dessa maneira, a transmissão de pacotes segue a mesma ordem idealizada pelo algoritmo BRR, porém sem o envio de bit-a-bit pelo enlace.

Seguindo este mesmo modelo é possível a atribuição de peso aos fluxos, ou seja, passa-se a supor uma transmissão de uma quantidade de bits proporcional ao peso do fluxo.

Controle Congestionamento

Entende-se por controle de congestionamento a capacidade que um equipamento de rede possui de gerenciar os momentos de congestionamento, de forma a minimizar a perda de pacotes o máximo possível. Os roteadores IP contam com os mecanismos disponibilizados pela diferentes implementações do protocolo TCP/IP.

Porém agregado ao controle do tamanho das janelas de transmissão, temos o controle das filas dos roteadores. Neste sentido, temos o mecanismo de controle WRED (*Weighted Random Early Detection*) que visa se antecipar às situações de congestionamento, então antes de um pacote adentrar na fila o mesmo será analisado para que se consiga perceber o fluxo, a precedência do pacote e a probabilidade de descarte associada à fila. Quanto maior for a prioridade dos pacotes, menores serão as suas chances de serem descartados.

No caso da necessidade de descarte, o mecanismo escolhe fluxos e descarta pacotes associados a estes fluxos até que a fila atinja um determinado patamar de ocupação.

2.2 Troca de Rótulos Multiprotocolar

Na tecnologia de “Troca de Rótulos Multiprotocolar (MPLS)” [RVC01], um pacote é encaminhado, no nível de rede, através do processamento padrão de cada pacote e, no nível de enlace, através da comutação por rótulos [MC01]. Para que este encaminhamento seja possível é necessário a presença dos roteadores de comutação por rótulos (LSR - *Label Switch Router*), que podem ser de Borda (*Edge LSR*) ou de Núcleo (*Core LSR*), os roteadores de Borda encontram-se na periferia da rede, enquanto que os de Núcleo se referem aos demais LSRs de uma rede MPLS.

Em um domínio MPLS, um caminho estabelecido para um dado pacote trafegar é denominado “Caminho Roteado por Rótulos - LSP (*Label Switched Path*)”. O LSP por natureza é unidirecional, um outro LSP deve ser criado para atender ao tráfego de retorno, ou seja, um LSP é uma sequência de LSRs ordenados onde o primeiro LSR é denominado LSR de ingresso (*Ingress LSR*) e o último LSR de Egresso (*Egress LSR*).

A rota estabelecida pelo LSP pode ser determinada por protocolos convencionais, por exemplo, OSFP (*Open Shortest Path First*) [Moy98], dessa maneira, o pacote percorre o mesmo trajeto dos datagramas convencionais. Uma outra forma de criar uma rota para um LSP é através do roteamento por restrições (CR - *Constraint-based Routing*), por exemplo, roteamento na origem ou roteamento com qualidade de serviço.

2.2.1 Classe Equivalente de Encaminhamento

Ao adentrar em um domínio MPLS, um datagrama é associado a uma Classe Equivalente de Encaminhamento (FEC- (*Forwarding Equivalence Class*)). Uma FEC com o qual o

pacote foi marcado é codificado com um valor de tamanho fixo conhecido como "rótulo".

Quando o pacote é enviado para o próximo nó o rótulo vai com ele. Nos nós subsequentes não existe uma outra análise do cabeçalho do pacote na camada de rede. Ao invés disso, o rótulo é usado como um índice para uma tabela que especifica o próximo nó e o próximo rótulo.

Durante o encaminhamento de um pacote com o uso da tecnologia MPLS, uma vez que ao pacote é associado a um FEC, outra análise por parte do próximo nó se faz desnecessária, todo o encaminhamento é feito por meio de rótulos. Uma série de vantagens são vistas a partir deste princípio.

- O encaminhamento MPLS pode ser feito por roteadores que são capazes de manusear rótulos, mas não são capazes de analisar cabeçalhos na camada de rede na velocidade adequada;
- Um pacote que acesse a rede por um roteador em particular pode ser rotulado de uma maneira diferente se o mesmo pacote entrasse na rede por um roteador diferente, sendo assim as decisões de encaminhamento que dependem do roteador de entrada podem ser tomadas com facilidade;
- As considerações que determinam como um pacote é associado a um FEC podem ter uma lógica complicada, sem impactar nos roteadores que fazem o encaminhamento dos pacotes rotulados;
- MPLS permite que o pacote siga um caminho pré-definido com facilidade, levando em consideração a engenharia de tráfego e o policiamento da rede. Já o encaminhamento tradicional, faz com que o caminho seja carregado pelo pacote;
- Alguns roteadores analisam o cabeçalho do pacote na camada de rede para não só decidir o próximo nó a ser seguido, mas também para determinar a precedência ou a classe de serviço. A tecnologia MPLS permite a inserção desta informação no rótulo, sendo o rótulo uma combinação de FEC com classe de serviço ou precedência.

2.2.2 Rótulos

Um rótulo é um identificador de tamanho fixo que é usado para identificar um FEC (Forwarding Equivalent Class) para o qual o pacote está associado. Na maioria dos casos, um pacote é associado a um FEC que, por sua vez, está se baseando no endereço de destino da camada de rede. Essa identificação é utilizada pelos LSRs de ingresso para decidir qual o encaminhamento dar aos datagramas que ingressaram em um domínio MPLS.

Nos LSRs de núcleo, os rótulos não são utilizados para o encaminhamento e sim para o controle do encaminhamento em dois momentos [MC01]:

1. quando do estabelecimento de um LSP, para se identificar, através da FEC, o próximo *hop* da rota;
2. quando as tabelas de roteamento são utilizadas, para certificar-se, através da FEC, se o próximo *hop* de uma rota estabelecida anteriormente para um LSP ainda permanece o mesmo (em caso de alteração, o LSP poderá ser desfeito ou re-roteado).

2.2.3 Agregação de Tráfego

Uma maneira de particionar o tráfego em FECs é pela criação de FECs separados para cada prefixo de endereço da tabela de encaminhamento. Contudo, dentro de um domínio MPLS, a consequência dessa atitude é um conjunto de FECs cujos tráfegos seguem os mesmos roteadores. Com isso, uma opção é a união deste tráfego em um único FEC por meio da agregação, reduzindo assim o número de rótulos para gerenciar e também a quantidade de distribuição de rótulos no controle de tráfego. Definindo assim uma agregação de tráfego.

2.2.4 Seleção de Rota

O protocolo MPLS possui duas opções de Seleção de Rota: (1) roteamento nó a nó e (2) rota explícita.

O roteamento nó a nó permite que cada nó, independentemente, escolha o próximo nó para cada FEC. Um LSP roteado nó a nó é um LSP cuja rota é selecionada usando o roteamento nó a nó.

No LSP roteado explicitamente, cada LSP não escolhe independentemente o próximo nó, ao invés, a maioria ou todos os LRSs do LSP foram selecionados com antecedência. A sequência de LSRs pode ser definida por configuração ou por um nó.

O roteamento explícito pode ser mais conveniente para o policiamento de roteamento assim como para o trabalho de engenharia de tráfego. No MPLS, a rota explícita precisa ter sido especificada no tempo que rótulos são associados. Isso faz com que o roteamento explícito seja mais eficiente do que o roteamento baseado no pacote IP.

2.3 Engenharia de Tráfego

Uma das mais significativas características do protocolo MPLS é a Engenharia de Tráfego [AMA⁺99].

O maior objetivo da Engenharia de Tráfego para a Internet é facilitar a eficiência e confiabilidade nas operações que a mesma suporta e, simultaneamente, otimizar a utilização de recursos e desempenho de tráfego. Com isso, a Engenharia de Tráfego se torna

uma função indispensável em muitos sistemas autônomos por causa do alto custo de manutenção da rede e natureza comercial e competitiva da Internet. Esses fatores enfatizam a necessidade de eficiência.

Em geral, esse conceito incorpora aplicações e princípios científicos de medidas, modelagem, caracterização e controle de tráfego na Internet para atingir objetivos específicos de desempenho. Os aspectos de Engenharia de Tráfego que estão relacionados ao protocolo MPLS são os de medida e controle.

Os objetivos chaves associados com Engenharia de Tráfego estão divididos entre:

- tráfego orientado
- recurso orientado

Os objetivos de desempenho para o tráfego orientado incluem aspectos que aperfeiçoam a QoS nos fluxos de tráfego. Em uma única classe, com modelo de serviço para a Internet de melhor-esforço, os objetivos de desempenho para o tráfego orientado são: minimizar a perda de pacotes, minimizar o atraso, maximizar a vazão e reforçar os acordos estabelecidos. Associado com os objetivos de desempenho do tráfego orientado se torna útil incluir diferenciação de serviços para a Internet.

Por outro lado, os objetivos de desempenho de recursos orientados incluem aspectos pertinentes à otimização da utilização de recursos, ou seja, com a garantia que sub-conjuntos de recursos da rede não sejam sobrecarregados e congestionados enquanto outros sub-conjuntos estão sub-utilizados, portanto a principal função da Engenharia de Tráfego é, eficientemente, gerenciar os recursos de banda.

Minimizar o congestionamento é um objetivo primário de ambos, ou seja, do tráfego e recurso orientado. Os problemas relacionados com congestionamento podem ter duração alta, sendo gerados por rajadas instantâneas.

O congestionamento tipicamente se manifesta de duas maneiras:

- Quando os recursos de rede são insuficientes ou inadequados para acomodar a carga recebida.
- Quando os fluxos não são mapeados de maneira eficiente de acordo com os recursos disponíveis gerando sobrecarga em algum destes, enquanto outros são sub-utilizados.

O problema de congestionamento pode ser resolvido por meio da ampliação da capacidade dos recursos ou pela aplicação de técnicas de controle de congestionamento ou mesmo pela combinação das medidas mencionadas anteriormente.

As técnicas clássicas de controle de congestionamento regulam a demanda do tráfego de forma a não sobrecarregar os recursos disponíveis, algumas técnicas são: controle de fluxo por janela, gerenciamento da fila de roteadores.

No geral, o congestionamento provocado pela má alocação de recursos pode ser amenizado pela aplicação de policiamentos de cargas. O objetivo de tais estratégias é o de minimizar o congestionamento ou alternativamente minimizar a utilização dos recursos, através da alocação eficiente dos mesmos.

No momento em que o congestionamento é diminuído através da eficiência da alocação de recursos, percebe-se uma menor perda de pacotes, menor atraso na rede e aumento na vazão dos dados, portanto um aumento na qualidade perceptível aos usuários finais. Sem esquecer que o balanceamento na carga também é uma fator importante no policiamento para a otimização do desempenho da rede.

2.4 MPLS e a Engenharia de Tráfego

O protocolo MPLS é apropriado à aplicação da Engenharia de Tráfego porque pode prover um número significativo de seus requisitos.

Antes de ir adiante com essa discussão, faz-se necessário definir “Tronco de Tráfego”. Tronco de Tráfego é uma agregação do fluxo de tráfego de uma mesma classe que é colocada dentro de um LSP, é uma representação abstrata de tráfego para com a qual algumas características específicas podem ser associadas [AMA⁺99].

Importante também se faz a distinção entre um “Tronco de Tráfego” e o LSP. Um LSP é uma identificação do caminho roteado por rótulos por onde o tráfego atravessa.

Retornando ao MPLS na implantação da Engenharia de Tráfego, pode-se dizer que essa característica está associada com os seguintes fatores:

1. Caminhos roteados por rótulos explícitos que podem ser criados administrativamente ou através de uma ação automática dos protocolos da camada inferior.
2. Os LSPs são gerenciados com eficiência.
3. Troncos de Tráfego podem ser mapeados em LSPs com facilidade.
4. Um conjunto de atributos podem ser associados aos “Troncos de Tráfego” considerando suas características comportamentais.
5. Um conjunto de atributos podem ser associados com os recursos que estão relacionados com os LSP e os Troncos de Tráfego.
6. MPLS permite tanto os tráfegos agregados quanto os desagregados.
7. Uma boa implementação do MPLS pode oferecer uma redução no cabeçalho para a Engenharia de Tráfego.

Em adição, através dos caminhos roteados por meio de rótulos explícitos, o MPLS permite uma aproximação da comutação por circuito que pode ser colocada como um modelo no roteamento na Internet.

2.5 MPLS e Serviços Diferenciados

Serviço Diferenciado é usado por alguns provedores de serviços para atender a redes que necessitem de múltiplas classes de serviços. Em algumas dessas redes, onde uma otimização apurada dos recursos de transmissão e melhora no desempenho e eficiência da rede é necessária, os mecanismos de engenharia de Tráfego podem ser complementares aos conceito de Serviços Diferenciados. Neste caso, o protocolo MPLS se apresenta como uma boa solução de integração de tecnologias.

Existem dois tipos de métodos de mapeamento de Serviços Diferenciados para o MPLS: L-LSP e E-LSP.

- **E-LSP (EXP-inferred-PSC LSPs)**

Um E-LSP pode carregar até 8 PHBs em um único LSP. Para este tipo de LSP, o campo EXP do cabeçalho do MPLS é usado pelo LSP para determinar o PHB a ser aplicado no pacote. Com este tipo de solução de mapeamento, como os tráfegos são roteados coletivamente, os mesmos compartilham o mesmo conjunto de limitações.

- **L-LSP (Label-Only-Inferred-PSC LSP)**

Levando-se em consideração este LSP, um único rótulo pode ser aplicado por LSP ou diversos PHBs com o mesmo regime de agendamento e diferentes prioridades de descarte, ou seja, por este LSP pode passar tráfego de AF11 e AF12. O PHB é determinado pelo rótulo ou pela combinação do rótulo com os bits do EXP.

Com esse modelo integrado existem também três adaptações no tocante ao processamento dos pacotes:

1. Além do *DSCP* do mecanismo de Serviços Diferenciados, um rótulo também é associado ao pacote no momento de entrada na rede, assim como o campo COS que se baseia na classe de serviço do pacote.
2. Durante o tráfego do pacote pela rede, a classe de serviço é definida pela interpretação do rótulo e/ou pelo campo COS, não mais pelo campo DS.
3. O desencapsulamento do pacote segue os preceitos definidos pela arquitetura dos Serviços Diferenciados.

2.6 Controle de admissão

Um algoritmo de controle de admissão determina se o novo fluxo de tráfego de tempo real pode ser admitido na rede sem prejudicar o desempenho garantido aos tráfegos já estabelecidos. Existem duas grandes classes de algoritmo de controle de admissão: controle de admissão determinístico e controle de admissão estatístico.

Para serviços de tempo real que precisam de ter um limite absoluto no controle de atraso de cada pacote, o modelo determinístico é usado. Em tais serviços determinísticos, o algoritmo de controle de admissão calcula o comportamento no pior caso de existência do fluxo corrente e do fluxo de entrada antes de decidir se o novo fluxo deve ser admitido. Este modelo subutiliza os recursos da rede, especialmente na ocorrência de um tráfego exponencial.

A maioria das aplicações como, por exemplo, fluxo de vídeo, precisa de garantias de desempenho muito rígidas e pode tolerar pequena violação nos limites de desempenho. Um controle de admissão estatístico pode ser usado neste caso, pois este algoritmo leva em consideração parâmetros estocásticos como entrada.

2.7 Políticas - gerenciamento de sistemas

A viabilização da implantação de QoS integrado com DiffServ requer um mecanismo de gerência de rede que simplifique a operação, levando em conta novos conceitos envolvidos na configuração da rede, e automatize algumas aplicações administrativas de rede.

A plataforma de políticas, padronizada pelo IETF, expõe o administrador da rede a ações de refinamento de processos, resolução de conflitos entre processos concorrentes, assim como estabilização de processos em um período do tempo.

2.7.1 Tipos de Políticas

As políticas podem também ser classificadas em dois grandes grupos [SL02]:

1. Políticas de Obrigação: são regras que são disparadas por eventos ou ações condicionadas, que podem ser usadas para definir reservas de recursos da rede, alterar estratégias de filas ou carregar um código para um roteador. Algumas dessas políticas são específicas ao usuário ou à aplicação.
2. Políticas de Autorização: são regras que definem serviços ou recursos que um usuário, uma regra ou um gerenciador de agentes podem acessar.

2.7.2 Gerenciamento de Políticas

O gerenciamento de políticas acontece em três etapas. Primeiramente, vem a etapa de "análise e diagnóstico" que atua na análise dos pontos positivos e negativos do domínio para, em seguida, ter início a etapa da "Escolha" que está relacionada com considerações de várias alternativas para decidir qual a melhor estratégia a ser seguida. Por último, apresenta-se a etapa da "Implementação" que aplica a estratégia definida durante a Escolha.

Levando-se em consideração estas características, o IETF/DMTF (*Internet Engineering Task Force/Distributed Management Task Force*) padronizou um modelo de administração baseado em políticas [RVKF99], conforme apresentado na Figura 2.1.

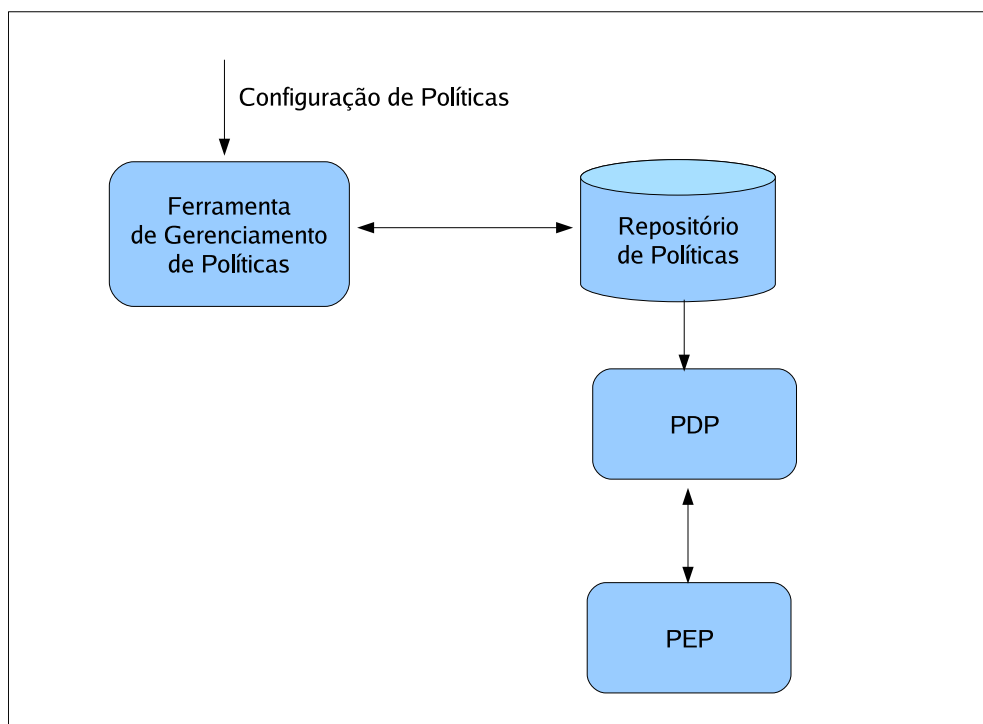


Figura 2.1: Arquitetura de Política do IETF/DMTF

A arquitetura de gerência de políticas apresentada na Figura 2.1 é composta por quatro elementos: Ferramenta de Gerenciamento de Políticas, Repositório de Políticas, PDP (*Policy Decision Point*) e PEP (*Policy Enforcement Point*).

A Ferramenta de Gerenciamento de Políticas é a interface com o administrador da rede, pela qual o mesmo é capaz de configurar as políticas de rede necessárias à gerência da rede.

Uma vez que as políticas foram validadas, as mesmas são armazenadas no Repositório de Políticas. As informações guardadas por este repositório obedecem ao formato estipulado pelo *Policy Framework Working Group* para que seja possível manter compatibilidade entre diferentes produtos [Dinesh 2002]. A comunicação com o repositório pode se dar por meio do LDAP (*Lightweight Directory Access Protocol*) [WHK06].

O PDP (*Policy Decision Point*) é responsável por atender aos pedidos de decisões impostos pelo PEP (*Policy Enforcement Point*). Uma vez de posse desse pedido, o PDP consulta o repositório de políticas e checa os recursos disponíveis para então enviar a decisão de operação para o PEP. De posse da política a ser aplicada, o PEP aplica a decisão tomada no nível dos pacotes que percorrem os elementos de rede supervisionados pelo PEP.

Em uma rede o PEP e o PDP podem estar em um único dispositivo ou em dispositivos diferentes. Caso eles se encontrem em localizações físicas distintas, a comunicação entre o PEP e o PDP pode se estabelecer por meio do protocolo COPS (*Common Open Policy Service*) [CSD⁺01] que se apresenta como uma possibilidade, assim como o protocolo SNMP (*Simple Network Management Protocol*) [HPW02].

Capítulo 3

Trabalhos Relacionados

Ao que se refere a políticas para uma plataforma MPLS-DS, que visa o oferecimento de serviços com a qualidade almejada pelo usuário final, o estudo dessa dissertação foi explorado em duas direções: uma relacionada à plataforma para prover este tipo de serviço fim-a-fim e outro relacionada ao gerenciamento das mudanças das condições de tráfego que podem ser variáveis e não previsíveis.

3.1 Gerenciamento de Tráfego

Com relação ao gerenciamento do comportamento da rede no que se refere ao tráfego em si, temos o LSP como uma entidade básica para adaptações em um rede baseada na tecnologia MPLS, ou seja, eficientes técnicas destinadas a esse gerenciamento estão associadas à conexão, desconexão e roteamento de LSPs, tendo como medida a capacidade da rede de adaptar suas configurações às mudanças de condições do tráfego.

No artigo [ASO05] são apresentados três algoritmos que baseados em medidas do estado corrente da rede conseguem adaptar as configurações da rede de modo a atender aos requisitos de tráfego pré-estabelecidos durante a admissão de tráfego.

O primeiro algoritmo mencionado, chamado de "*MDP-Based Optimal LSP Setup*", aparece como um estudo de caso para a situação na qual existe um pedido de alocação de banda entre dois nós que não estão conectados via um LSP. A política discorre na maneira como será realizada a alocação do LSP, chegando a uma conclusão que a capacidade de banda do novo LSP deve ser igual ao valor nominal dos pedidos de banda, devendo ser mantida constante até o próximo pedido de redimensionamento. Um problema encontrado para este algoritmo é o de que a banda não pode ser completamente utilizada durante um espaço de tempo muito significativo, devido às mudanças nos perfis de transmissão dos fluxos correntes.

O segundo algoritmo, proposto por [ASO05], se preocupa com o dimensionamento do

tráfego de uma classe de tráfego em um LSP previamente estabelecido entre dois nós. Este algoritmo pode ser usado para completar as políticas que decidem sobre a criação, redimensionamento e desconexão de um LSP. Este algoritmo fez uso do filtro de Kalman, devido às impurezas de dados relacionadas a esse assunto.

Já o terceiro algoritmo volta-se para o problema de roteamento de LSPs. Este algoritmo introduz um filtro que diminui a chance de erro durante o roteamento.

Outras duas políticas foram propostas no artigo [BQ01], a primeira delas é aplicada para um pedido de banda menor do que 512 Kbits/s, para o qual o tráfego é mapeado para um LSP, previamente, criado. Para pedidos de banda acima de 512 Kbits/s, um LSP dedicado é criado a fim de mapear o novo tráfego entrante.

Para este artigo, as conclusões se ativeram às restrições oferecidas pela plataforma proposta pelo IETF. Segundo o artigo, na plataforma em questão, não é possível a detecção de conflitos, assim como as definições de ações e condições não são genéricas, levando a uma especialização sem base para futuras referências. Porém para essas conclusões e propostas de políticas não foram apresentados dados para verificação e validação.

Outro problema relevante nesta área se refere à necessidade de re-roteamento. Seguindo esta linha de estudo, o artigo [dOSAU04] preocupa-se em propor políticas de prevenção que tentam assegurar que os LSPs de alta prioridade possam ser roteados através de caminhos que suportem serviços diferenciados.

V-PREPT é uma política de prévia apropriação que combina os três principais critérios de otimização: número de LSPs, prioridade dos LSPs a serem pré-definidos e quantidade de banda a ser pré-estimada. Outra política de antecipação explorada neste artigo foi a Adpat-V-PRET que seleciona os LSPs de baixa prioridade a fim de reduzir sua capacidade máxima de alocação de banda a fim de abrigar no enlace um LSP de maior prioridade.

A eficácia destas políticas foi validada através de simulações e resultados dos experimentos, pelos quais conclui-se que as políticas de antecipação propostas no artigo em questão mostraram um desempenho melhor do que as políticas não consideradas de antecipação em um ambiente de serviços diferenciados.

Outro ponto que pode ser ressaltado é a relação entre prévia apropriação de recursos e a seleção de um caminho. A seleção de um caminho na plataforma MPLS-DS, na sua maioria, está baseada em CBR - "*Constraint-Based Routing*", que se mostra capaz de encontrar vários caminhos que satisfaçam as mesmas considerações, além de escolher um caminho preferencial a partir da seleção executada, previamente, e, eventualmente, otimizando a utilização de recursos.

Enquanto que, durante a computação por uma rota para um LSP, muitas das vezes, considera-se, apenas a alocação de recursos no nível de maior prioridade do LSP e informações sobre outros LSPs com um nível de prioridade mais baixo são ignorados, com isso pode ser que aconteça de um novo LSP se apropriar de um conjunto de LSPs de

menor prioridade de um determinado caminho, sendo que um outro caminho não sofre essa apropriação de recursos.

Levando em consideração os aspectos mencionados acima, o artigo [KLHm04] propõe um algoritmo de seleção de caminho que leva em consideração a apropriação prévia, ou seja, considera a reserva de recursos por parte dos LSPs de menor prioridade na seleção de um caminho, na tentativa de evitar reroteamento e aumentar a estabilidade da rede.

As simulações foram capazes de provar que para determinadas situações, o algoritmo apresentou uma diminuição de reroteamento e de balanceamento de carga, quando comparado com os algoritmos de seleção de caminho pela menor distância [MSZ96] e pela apropriação prévia e vetorial de largura de banda [SSJ02].

Já o artigo [CHL⁺03] propõe uma política que evita a apropriação prévia por parte dos fluxos de maior prioridade e, ainda assim, balanceamento de carga para uma rede DS-TE. O algoritmo calcula a probabilidade de apropriação prévia do LSP por outro de maior prioridade, dessa maneira, evita a desapropriação do LSP.

Enquanto que a dissertação [Gui01] se preocupa em validar políticas onde se estipula critérios para a criação de LSPs com novas rotas (re-roteamento) ou mesmo políticas que se baseiam na configuração dos pesos das filas que estão associadas a LSR segundo a classe de serviços e o mapeamento das classes de serviços em um LSP.

De acordo com a dissertação mencionada acima, a política de criação de LSPs é ativada quando o atraso relacionado a uma classe de serviço for ultrapassado. Enquanto que, a política de configuração das filas é ativada caso a porcentagem do descarte dos pacotes pertencentes a uma mesma classe de serviço seja excedida. Depois de três tentativas, caso a situação não se estabilize, ocorre um re-roteamento de tráfego, através da criação de um novo LSP. A política de mapeamento das classes de serviço em um LSP é realizada implicitamente quando um novo LSP é criado.

As políticas apresentadas acima foram consideradas satisfatórias para os cenários de simulação aos quais foram expostas.

A corrente dissertação leva em consideração os artigos e a dissertação citados acima no sentido em que propõe novas políticas de antecipação, bem como políticas de atuação. Porém, diferentemente, da maioria das políticas mencionadas acima, as políticas que serão expostas levaram em conta não só a largura de banda como o nível de descarte de pacotes. Estas políticas serão explicadas, com maiores detalhes, no Capítulo 4.

3.2 Plataformas

No artigo [ASO05] temos a apresentação da plataforma TEAM que foi desenvolvida para o gerenciamento de um domínio DiffServ/MPLS, sendo responsável pela alocação de banda e roteamento do tráfego entrante e pela tomada de decisões dinâmicas que afetam a

configuração da rede. Esta plataforma foi usada como base para a validação dos algoritmos apresentados pelo artigo [ASO05], mostrando-se satisfatória para os propósitos do artigo.

Para o artigo [BQ01], um protótipo do PDP foi desenvolvido, baseado na proposta exposta pelo IETF *Policy Framework*, para o gerenciamento da rede MPLS. O mesmo artigo não apresenta dados de verificação e validação, não sendo possível uma avaliação do PDP desenvolvido.

Para a validação do algoritmo citado em [CHL⁺03], temos uma plataforma desenvolvida seguindo os critérios do NS-2 [NS01]. Os resultados apresentados pela plataforma mostraram que a política proposta aplicada ao modelo tradicional de CBR melhora o funcionamento do mesmo, mesmo percebendo que os fluxos de maiores prioridades não trafegaram pelo menor caminho.

Ao se falar na plataforma NS-2 temos dois trabalhos que se basearam nesta plataforma e que tiveram influência direta no desenvolvimento desta dissertação.

O primeiro trabalho é o [Gui01] que utiliza uma plataforma MPLS-DS para validar políticas de configuração.

A plataforma MPLS-DS é composta por cinco elementos: um *Bandwidth Broker*, a rede MPLS-DS proporcionada pelo NS-2, um mecanismo de mapeamento de classes de serviços, a arquitetura de políticas, baseado no modelo de políticas de três camadas (*Three -Tier Policy Model*) [RVKF99] e os contratos SLs.

O segundo trabalho [Ped04] propõe um modelo dinâmico e automático de negociações de serviços para a Internet, simplificando o processo de administração e configuração e garantindo a realização de transações fim-a-fim com requisitos de QoS.

Para tal, utilizou-se a plataforma proposta por [Gui01] para implementar um protocolo baseado nos protocolos COPS, BGRPP e SIBBS, assim como foram disponibilizadas políticas de negociação de recursos, adotando a estratégia de negociação atacado/varejo.

Através de simulação foi comprovada a eficácia do modelo de negociação.

Esta dissertação faz uso da plataforma MPLS-DS proposta por [Gui01], contando com o modelo de negociação proposto por [Ped04], para validar políticas de configuração de rede, que levam em consideração fatores como descarte de pacotes e o consumo de banda do tráfego no *backbone*.

Para a implementação das novas políticas foi necessária uma melhoria no sistema de Engenharia de Tráfego da plataforma corrente, de maneira a possibilitar um acompanhamento mais detalhado dos fluxos de tráfegos de um LSP, levando não só em consideração a classe de serviço do mesmo, como também o destino do mesmo. No próximo capítulo, esse aspecto da plataforma será explorado.

Capítulo 4

Plataforma Proposta

Este capítulo apresenta a plataforma que foi usada para simular uma rede MPLS-DS, que provê QoS fim-a-fim, para a validação de uma proposta de melhora no mecanismo de “Engenharia de Tráfego” existente no simulador de rede MPLS. Assim como, para a validação de um conjunto de políticas de tráfego que visam o dinamismo e a garantia de cumprimento de contratos, devidamente, pré-estabelecidos com os clientes da rede no momento da entrada do tráfego associado no domínio.

A plataforma pode ser dividida em dois conceitos macros: o “*BB - Bandwidth Broker*” e o “*NS - Network Simulator*”. O BB está associado ao gerenciamento de recursos da rede mediante o tráfego exposto, enquanto que o NS gerencia e configura os elementos de rede.

4.1 Network Simulator

Primeiramente, faz-se necessário discorrer sobre o que vem a ser o “NS - Network Simulator”.

O NS é uma ferramenta poderosa, desenvolvida pela universidade de Berkeley, que dirigida por eventos simula vários tipos de redes IP. Esta ferramenta implementa protocolos de redes tais como o TCP e o UDP, assim como aplicações de FTP(*File Transfer Protocol*), Telnet, VBR(*Variable Bit Rate*), CBR(*Constant Bit Rate*), dentre outras, possuindo também mecanismos de gerência de filas de roteadores como *Drop Tail*, RED(*Random Early Detection*) e CBQ(*Class-Based Queuing*), além de algoritmos de roteamento como o OSPF dentre outros. Na plataforma, em questão, utiliza-se a versão 2 do mesmo.

Como mostra a Figura 4.1, numa visão simplificada do usuário, o NS é um interpretador de *script* Tcl orientado a objeto (OTCL - *Object-Oriented Tool Control Language*) que tem como suporte uma biblioteca que contém objetos de agendamento de evento, objetos de componentes de rede e módulos de ajuda de configuração de rede.

Em outras palavras, para configurar e rodar um simulador de rede, o usuário deve:

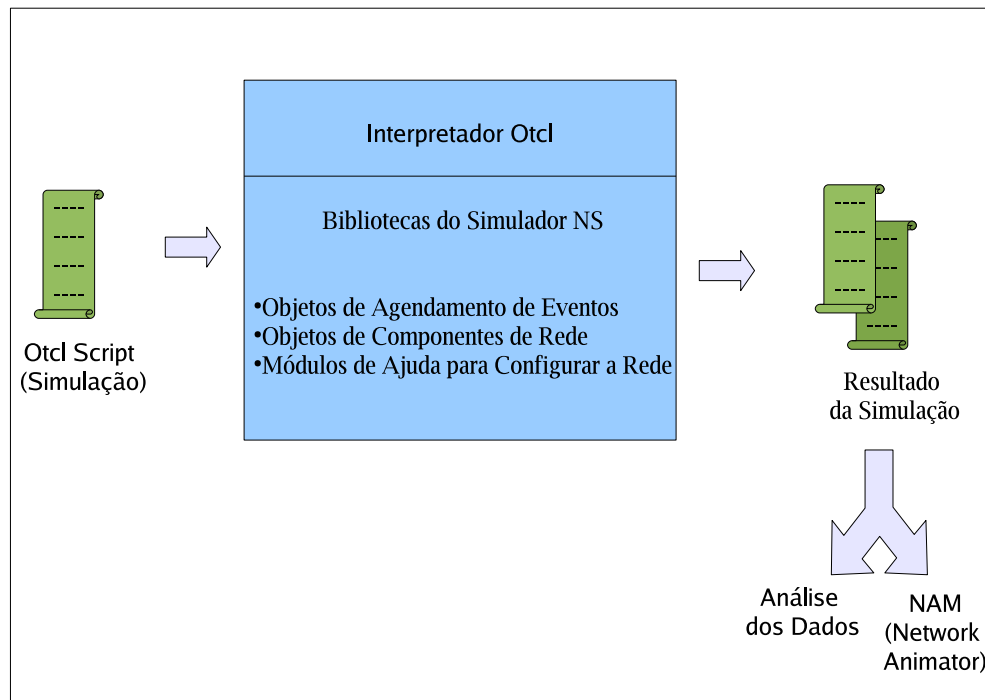


Figura 4.1: Visão Simplificada do NS

- escrever um script OTcl que inicia um agendador de evento;
- configurar a topologia da rede utilizando os objetos de rede e módulos de junções;
- informar as fontes geradoras de tráfego quando iniciar e parar a transmissão de dados através do agendador de eventos.

O módulo de junção é utilizado para configurar a rede no sentido de que caminhos de dados entre os objetos de rede são dados pela configuração do ponteiro do "vizinho" de um objeto para o endereço de um objeto apropriado. Quando um usuário quer construir um novo objeto de rede, este pode facilmente fazê-lo construindo um novo objeto ou utilizando um objeto de uma biblioteca, e ligando o caminho de dados pelo objeto.

Por razões de eficiência, o NS é implementado em C++ e usa o OTCL como uma interface de comandos e configuração. Objetos de componentes de rede como os módulos MPLS, DS e o MPLS-DS são escritos e compilados utilizando C++, sendo disponibilizados para o interpretador de OTCL por uma ligação que cria uma correspondência do objeto OTCL para cada um dos objetos C++.

As funções de controle e as variáveis de configuração especificadas pelo objeto C++

agem como funções de membros e variáveis membros de objetos OTCL correspondentes. Desta forma, os controles dos objetos C++ são configurados via OTCL, tornando possível adicionar funções membros e variáveis para os objetos C++ ligados aos objetos OTCL. O objeto em C++ que não é necessário ser controlado na simulação ou internamente utilizado por um outro objeto, não precisa ser ligado ao OTCL.

Nesta dissertação foi usado o módulo MPLS-DS proposto por [Gui01] com algumas adaptações, assim como foi acrescentando o novo agente TCP/SinkMonitor [Mon] na plataforma.

Primeiramente, as adaptações do módulo MPLS-DS dizem respeito ao acréscimo de uma nova tabela *LSPFlow* no mecanismo de “Engenharia de Tráfego” existente. É responsabilidade da tabela *LSPFlow* manter a relação entre LSP, classe de serviço e destino de um pacote de entrada, possibilitando a criação de políticas de prevenção.

Em segundo lugar, temos um aumento no monitoramento dos fluxos correntes em um LSP através do acompanhamento da ocupação de banda do mesmo.

4.1.1 MPLS-DS

Um nó MPLS-DS é capaz de encaminhar um tráfego por meio do protocolo MPLS e ainda diferenciar os diversos tipos de serviços. Este nó foi composto pela integração dos módulos DS-NORTEL[Pie00] e MPLS, que atende pela sigla MNS[AC01].

DS-Nortel

Por meio do pacote DS-Nortel, as filas de borda e de núcleo, encontram-se disponíveis. As filas de borda realizam a classificação, marcação, policiamento, adequação e encaminhamento de pacotes, enquanto que as filas de núcleo encaminham os pacotes com base no PHB da classe ao qual o pacote pertence.

Cada fila física pode ter três filas virtuais para a diferenciação de serviços, sendo quatro o número máximo de filas físicas.

MNS

O pacote MNS oferece ao módulo MPLS-DS as características de roteamento por rótulo, o manuseio de mensagens LDP (*Label Distribution Protocol*), assim como o de mensagens CR-LDP.

Para a caracterização de um nó MNS, faz-se necessário agentes e classificadores. O “agente” é um objeto de envio e recepção de mensagens pertencentes a um protocolo, enquanto que o “classificador” classifica os pacotes de acordo com as normas

pré-estabelecidas. Dessa maneira, para um nó MNS, existem os objetos "Classificador MPLS" e "Agente LDP" para diferenciar de um nó IP comum do NS.

O "Classificador MPLS" verifica se o pacote está rotulado ou não. Caso o pacote já esteja rotulado, o mesmo é, devidamente, encaminhado para o próximo nó. Caso o pacote não esteja rotulado, na existência de um LSP para o seu endereço de destino, o mesmo é manuseado como um pacote rotulado, caso contrário, será enviado para o "Classificador de Endereço" do NS.

Um nó MNS tem três tabelas para gerenciar informações relacionadas com o LSP:

- PFT - *Partial Forwarding Table*;
- LIB - *Label Information Base*;
- ERB - *Explicit Routing Information Base*.

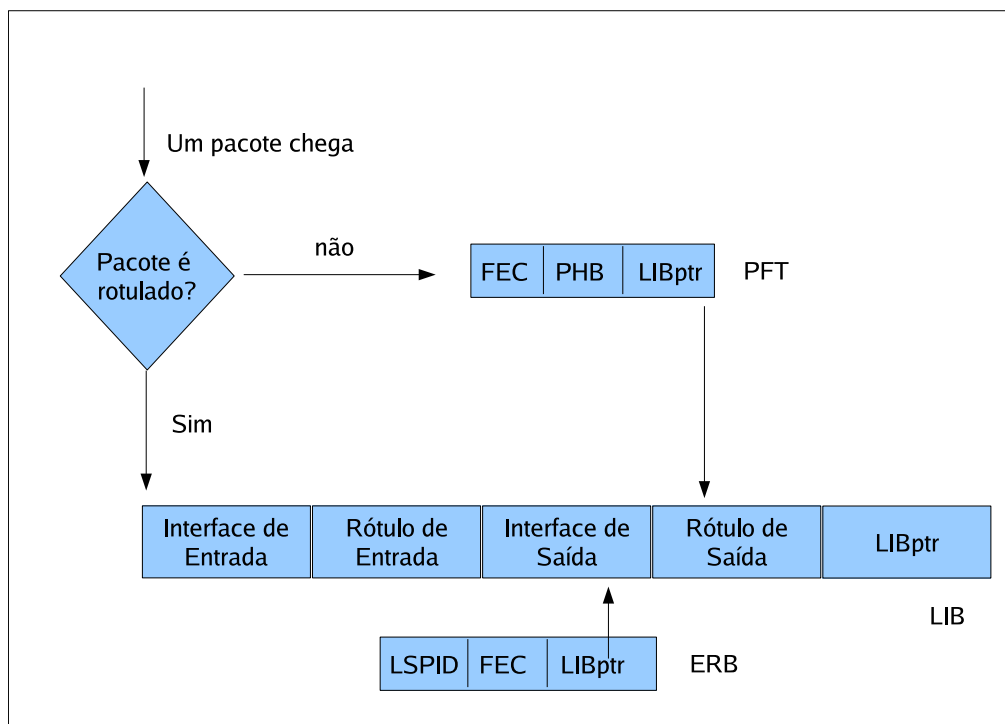


Figura 4.2: Estrutura de Tabelas do MNS

A tabela PFT é um sub-conjunto da tabela de encaminhamento que contém o FEC, onde o PHB é um apontador para a tabela LIB que por sua vez contém informações sobre

os LSPs que estão estabilizados. Já a tabela ERB contém informações sobre o ER-LSP, ver Figura 4.1.1.

A procura na tabela PFT/LIB se inicia quando um nó MPLS recebe um pacote. No caso de um pacote não rotulado, o classificador MPLS identifica a localização na tabela PFT pelo FEC, se o apontador para a tabela LIB do mesmo indica valor nulo, então o pacote será encaminhado para o "Classificador de Endereços", caso contrário, o pacote é encaminhado para o próximo nó via um roteamento do tipo L2, indicado pelo campo "*outgoing-interface*" da tabela LIB.

Se acontecer do pacote estar rotulado, a tabela LIB é diretamente consultada e o rótulo atual do pacote é trocado pelo valor do campo "*outgoing-interface*" e, em seguida, direcionado ao próximo nó pelo roteamento do tipo L2.

A tabela ERB não é consultada para o encaminhamento do pacote. Caso seja necessário mapear um novo fluxo em um ER-LSP, previamente criado e estabilizado, uma nova entrada na tabela PFT com o mesmo apontador para a LIB e ERB deve ser acrescentado.

Porém de posse dessa três tabelas não se faz possível gerenciar tráfegos com o mesmo PHB dentro de um único LSP. Neste cenário, as políticas de reconfiguração de filas ou mesmo de rearranjo de peso de filas não são eficientes, pois as mesmas só levam em consideração o parâmetro PHB.

Pensando nesta direção, este trabalho criou uma nova tabela LFT (*LSP Flow Table*) pela qual é possível diferenciar os fluxos em um LSP não só pelo PHB, mas também o destino dos pacotes.

Dessa maneira, para o cenário em que um LSP possui todos os seus fluxos pertencentes a uma única classe de serviço, alguma ação pode ser tomada no sentido de garantir o QoS de fim-a-fim. Esta nova tabela e suas implicações na plataforma serão explicadas com mais detalhes na próxima sub-seção "LFT - LSP Flow Table".

LFT - LSP Flow Table

Para que se fizesse possível um acompanhamento mais detalhado dos fluxos de tráfego que ocupam um LSP, criou-se a nova tabela LFT que diferencia as classes de serviço existentes no LSP pelo destino final do pacote, melhorando a Engenharia de Tráfego do sistema.

Para o correto gerenciamento dessa tabela a classe MPLDSAddressClassifier [Gui01] foi alterada com a inclusão de novas estruturas de dados, bem como criação de novos métodos e alteração de outros pré-existentes. É importante ressaltar que a introdução da tabela FLT não altera o encaminhamento do tráfego.

A Figura 4.1.1 mostra os métodos alterados e novos, bem como as novas estruturas de dados auxiliares criadas para a classe MPLDSAddressClassifier.

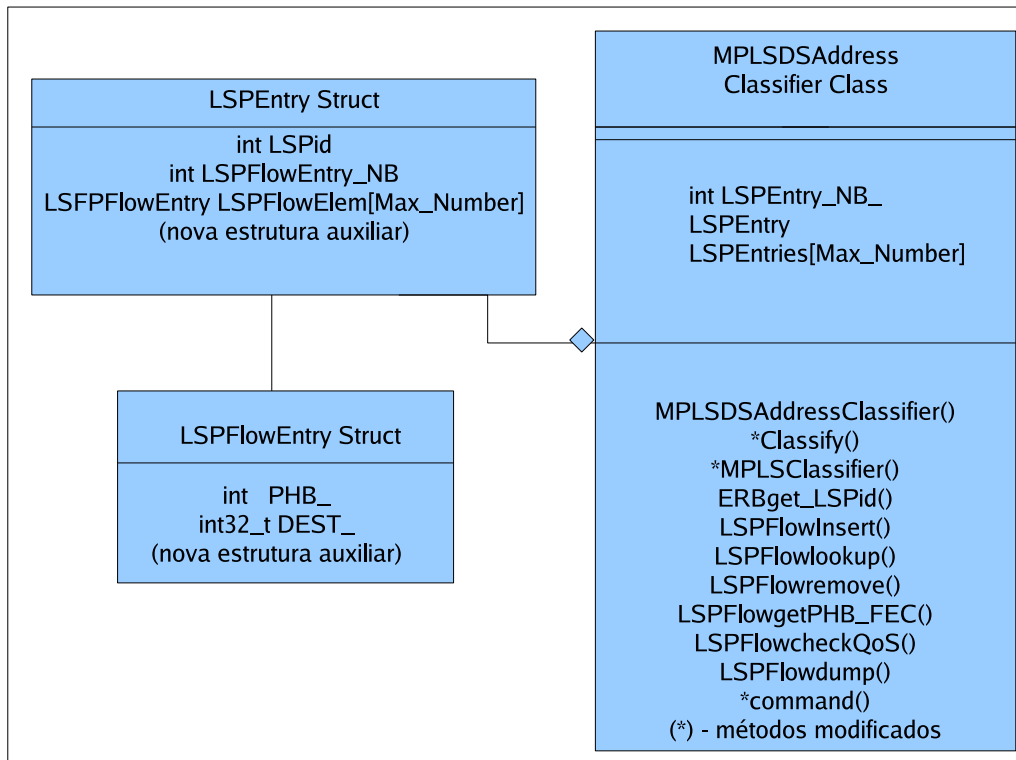


Figura 4.3: Classe MPLSDSAddressClassifier

Como pode-se verificar na Figura 4.1.1 foi acrescentado à classe MPLSDSAddressClassifier os atributos:

- `LSPEntry_NB_` - Guarda os números de LSPs criados no MPLS-DS de ingresso.
- `LSPEntry LSPEntries[Max_Number]` - Defini um vetor que irá guardar dados de um LSP. O número máximo de LSPs que podem trafegar no sistema é configurável e dependente do desempenho da plataforma. Para os experimentos do Capítulo 5, estipulou-se o valor de 10 LSPs.

A estrutura `LSPEntry` é definida pela identificação do LSP e pela estrutura `LSPFlowEntry` que guarda a identificação dos PHBs e, respectivos, destinos que trafegam pelo LSP. Neste caso pode ter um máximo de 8 `LSPEntry` por LSP, pois os LSP, em questão, são do tipo ELSP.

Já os métodos listados na Figura 4.1.1 podem ser descritos da seguinte maneira:

- `MPLSAddressClassifier` - Construtor. Inicializa as variáveis de controle.

- Classify - Verifica para onde deve encaminhar um pacote IP de acordo com o rótulo. A este método foi acrescentada a habilidade de inserir na tabela LFT os dados de PHB e destino para um pacote, contando com a ajuda da leitura das tabelas MPLS PFT e ERB.
- LSPFlowinsert - insere uma nova entrada na tabela LFT.
- LSPFlowremove - remove uma entrada na tabela LFT.
- ERBget_LSPid - Através da tabela MPLS ERB chega-se ao LSP ao qual o pacote está associado. A entrada deste método é o apontador da tabela LIB que se obtém por meio da tabela MPLS PFT.
- LSPFlowlocate - Retorna a posição do LSP no vetor LSPEntries.
- LSPFlowlookup - Verifica se a tabela LFT já possui entrada para o LSP e, respectivos, PHB e destino.
- LSPFlowgetPHB_DEST - Retorna o valor específico do PHB e destino associado de acordo com a identificação do LSP e posição no vetor LSPFlowElem.
- LSPFlowcheckQoS - Verifica se todos os fluxos de um dado LSP possuem a mesma classe de serviço.
- LSPFlowdump - Imprimir a tabela LFT.
- command - Permitir ao monitor de LSP acesso aos dados da tabela LFT.

Em termos de escalabilidade da rede como um todo, o impacto da tabela LFT pode ser considerado como de baixo para o sistema, pois a mesma só é utilizada nos roteadores de borda - *Edge* LSRs, não interfere no encaminhamento do datagrama e solicita uma baixa demanda na capacidade de processamento, pois a chave primária para a tabela é o LSPid que sendo do tipo ELSP pode ter no máximo 8 destinos.

Quanto ao fator armazenamento, o mesmo está diretamente relacionado com o número de LSPs que podem ser criados no sistema, o que pode ser configurado de maneira que não afete o encaminhamento dos pacotes.

Em resumo, sem muito impacto computacional, a perspectiva de tabelas do MNS foi alterada, segundo a Figura 4.1.1, onde através da tabela LFT é possível capturar uma classe de serviço e o destino de um fluxo dentro de um LSP.

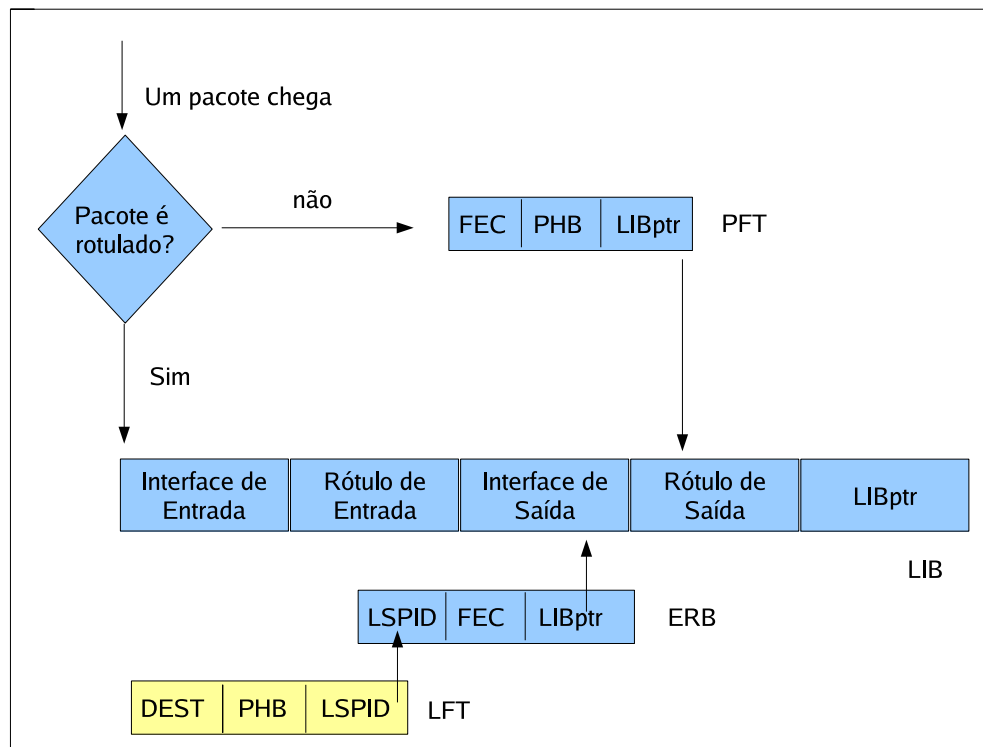


Figura 4.4: Estrutura de Tabelas do MNS incluindo a Tabela LFT

4.1.2 Integração DS-NORTEL & MNS

Com a integração, entre os módulos MNS e DS-Nortel, todos os pacotes são rotulados na entrada do domínio e aos roteadores LSR-DiffServ cabe a função de interpretar o campo EXP, assim como associar o destino do pacote IP à classe de serviço, como foi dito, anteriormente.

Através da integração também será possível a criação de dois tipos de LSPs: E-LSP e L-LSP.

4.1.3 Agente TCP/SinkMonitor

Na linguagem do NS, um agente é uma classe escrita em C++ que representa um ponto de construção ou consumo de pacotes.

As estatísticas mantidas por este agente são:

- *recvbytes_* : número total de *bytes* enviados pela camada de aplicações.
- *nlost_* : número de pacotes perdidos. Este número é calculado pelo número de

seqüência no cabeçalho dos pacotes TCP. Se o pacote recebido tem o número de seqüência maior do que o esperado, isso significa que alguns pacotes foram perdidos.

- *ndup_* : número de pacotes duplicados. Este evento pode ocorrer desde que alguns *acks* sejam perdidos. Este valor é calculado pelo número de seqüência no cabeçalho do pacote TCP. Se o pacote recebido tem o número de seqüência menor do que o esperado, isso significa que o fonte re-enviou alguns pacotes devido à perda ou estouro do tempo de confirmação de recebimento do pacote por parte do receptor.
- *expected_* : número de seqüência do próximo pacote a ser esperado.
- *npkts_* : número total de pacotes recebidos, incluindo os duplicados.

4.2 Arquitetura do Bandwidth Broker

A implementação do BB utilizada nesta plataforma segue o modelo Cliente/Servidor. Os servidores podem ser enumerados como de Negociação, de Recursos e de Política, que são responsáveis pela tomada de decisões e armazenamento de informações. Com relação aos clientes, temos os objetos de configuração, aplicação e gerenciamento. A Figura 4.2 demonstra o modelo descrito.

O BB foi escrito em iTCL[ITC] que, em linhas gerais, é uma extensão orientada a objetos da linguagem TCL. A comunicação entre o BB e a rede MPLS-DS se estabelece através da arquitetura CORBA(*Common Object Request Broker Architecture*)[Gro], possibilitando ao BB a migração para um ambiente real, exercendo a função de gerenciador dos recursos da rede. Através do pacote COMBAT[TCL] e suas interfaces DII(*Dynamic Invocation Interface*) e DSI (*Dynamic Skeleton Interface*) é feito o mapeamento do iTCL para o CORBA. O COMBAT permite a criação de *scripts* tanto do lado servidor quanto do lado cliente.

A sinalização entre as mensagens trocadas pelos servidores e clientes segue o protocolo COPS ou COPS-SLA. A emulação das mensagens de inicialização do protocolo COPS é utilizada entre os elementos Objeto de Configuração e Gerente de Políticas.

O protocolo COPS-SLA [Ped04] é executado para a comunicação entre o Gerente de Negociação e o Objeto de Aplicação. O protocolo COPS-SLA possui duas novas entidades o SLA-PDP e o SLA-PEP que se comportam, basicamente, como o PDP e o PEP do protocolo COPS, a diferença entre os protocolos reside no fato de que o COPS-SLA não só se orienta pelas políticas de negociação estipuladas pelo domínio, mas também nos contratos de serviços (SLAs) que foram estabelecidos junto aos clientes para a tomada de decisões sobre a concessão de recursos.

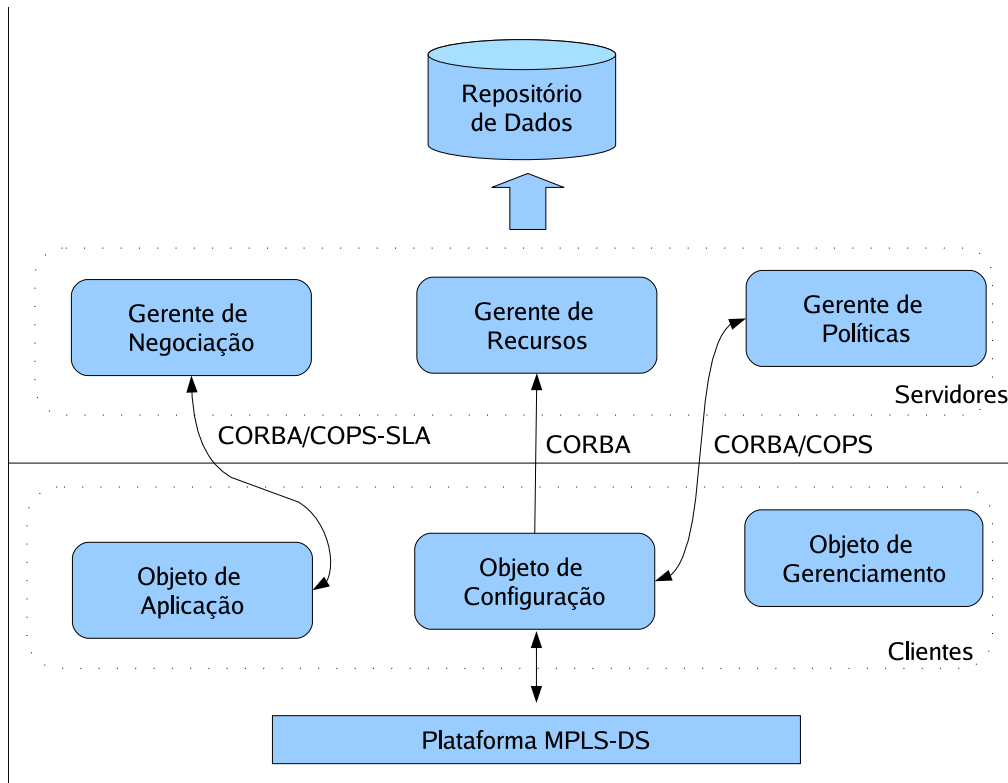


Figura 4.5: Arquitetura do Bandwidth Broker

Outra diferença marcante está entre o PEP e o PEP-SLA. O PEP-SLA não é um aplicador de políticas, como o PEP nos elementos de rede, e sim um configurador de contratos para o gerente de negociação. Quando existe um contrato que foi renegociado, o gerente de negociação utiliza o PEP para a reconfiguração dos elementos de rede.

4.2.1 Gerente de Recursos

O gerente de recursos com base no conhecimento dos contratos em andamento no domínio e estado dos enlace, através do PEP, pode monitorar o cumprimento dos contratos. Em caso de quebra de contrato, o gerenciador de Políticas é notificado para que tome as medidas necessárias para que o QoS não seja afetado.

Além dessa funcionalidade, o gerente de recursos através de sua interação com o PEP pode oferecer e garantir recursos no momento da negociação.

4.2.2 Gerente de Políticas

O gerente de políticas representa o PDP que tem como função tomar decisões que garantam o QoS definido no contrato do cliente por meio de políticas armazenadas no repositório de políticas.

Esse mecanismo se solidifica por meio do COPS que estabelece uma ponte entre o PDP e o PEP.

4.2.3 Gerente de negociação

O gerente de negociação é responsável pela negociação e gerenciamento de contratos de serviços, através de políticas de negociação, proporcionando dinamismo e automatização ao sistema.

A fim de garantir QoS fim-a-fim, a negociação de contrato de serviços pode ser intra ou inter domínios. A negociação intra domínio se realiza entre o cliente e o gerente de domínio, enquanto que a negociação inter domínios acontece entre os gerentes de domínios.

A estratégia de negociação empregada pela plataforma segue o modelo atacado/varejo, caracterizando como varejo a negociação entre domínios vizinhos de forma contínua até que se alcance o domínio vizinho ao domínio de destino, através desse modelo de negociação, conhecido como modelo cascata, criam-se "túneis" de QoS entre os domínios e uma redução significativa no número de mensagens de sinalização.

No modelo de negociação por atacado, o gerente de negociação oferece ao cliente seus recursos locais, adicionando os recursos adquiridos pela negociação por varejo.

4.2.4 Objeto de Configuração

O Objeto de Configuração é representado pelo PEP, o mesmo está presente em cada roteador do domínio, sendo o ponto de intersecção entre os planos gerencial e operacional da plataforma. O PEP faz parte da plataforma MPLS-DS.

Ao objeto de configuração compete as funções de monitoramento de tráfego e configurador de roteadores, mediante instruções provenientes do Gerente de Políticas.

4.2.5 Monitoramento de Tráfego Simulado

O nó MPLS-DS disponibiliza o monitoramento de tráfego de um LSP. O monitoramento é implementado em OTCL.

Os parâmetros que podem ser monitorados na plataforma são: descarte e atraso de pacotes, assim como a largura de banda de um LSP.

Para a medição do nível de descarte na rede, cada roteador apresenta mecanismos de controle de pacotes recebidos, porém que não foram encaminhados, através do monitoramento das filas.

Enquanto que para a medição do atraso, verificou-se que o valor da medida provinha do monitoramento das filas dos roteadores, de forma a medir o tempo de chegada do pacote na fila e o tempo de saída da fila, e não o tempo que o pacote realmente leva para trafegar pela rede. Dessa maneira, para esta dissertação, os valores de atraso não foram considerados como suficientes para o acionamento de políticas.

Outro fator que contribui para a não utilização do atraso como acionador de políticas é a inexistência de um serviços de temporização para redes reais o que impede a medição do tempo que um pacote leva da sua origem até o seu destino.

Através das medidas de descarte e ocupação de banda de um LSP foram criadas novas políticas de atuação para a plataforma em questão. Dentro do conjunto de políticas de atuação, as políticas de prevenção tomaram forma por meio do acompanhamento mais detalhado do tráfego proveniente da tabela LFT, vide sub-seção 4.1.1.

4.2.6 Objeto de Aplicação

O Objeto de Aplicação é representado pela Figura do PEP-SLA na arquitetura de políticas, encontra-se presente nos roteadores de borda ou em um gerente de domínio vizinho.

A função do mesmo é servir como interface a um cliente que deseje fazer reserva de recursos.

4.2.7 Objeto de Gerenciamento

O objeto de gerenciamento é a interface entre o administrador da rede e os gerentes de serviços do BB.

Através dessa interface, o administrador é capaz de executar instruções de configuração sobre os três módulos gerenciais.

4.2.8 Repositório de Dados e Contratos de Serviços

O repositório de dados e serviços é acessado pelo gerente de negociação que interpreta os dados e os armazena para que sejam, devidamente, utilizados pelo gerente de recursos. A organização dos dados se dá na forma de contratos de serviços no formato de arquivos textos.

Os contratos de serviços podem ser especificados pelo cliente ou pelo provedor de serviços, o formato adotado nesta plataforma está dividido em três partes: unidade comum, unidade de topologia e unidade de QoS.

Unidade Comum

A unidade comum é composta de informações sobre a identificação e controle de contrato.

- Formato de Contrato - tipo de contrato utilizado na negociação.
- Especificação de Serviços - identifica os serviços negociados no contrato.
- Identificação de Provedor - identifica o domínio que esta cedendo recursos mediante o contrato de negociação;
- Identificação de Cliente - identifica o solicitante.
- Tipo de Cliente - classifica o solicitante como gerente do domínio ou cliente local.
- Data de criação do contrato.
- Data de Validade.

Unidade de Topologia

Esta unidade se encarrega de carregar informações sobre os nós envolvidos no trajeto do pacote de fim a fim. O cliente deve preencher os dados do nó de origem e o provedor deve preencher os dados sobre o nó de destino, sendo assim o cliente não precisa conhecer a topologia do provedor, e vice-versa.

Unidade de QoS

Esta unidade carrega consigo todas as informações necessárias referente ao domínio e ao serviço. Para uma melhor organização desta unidade, a mesma é composta por uma lista de sub-unidades de QoS que estão listadas abaixo:

- Sub-unidade de Escopo - identifica os nó de fim-a-fim;
- Sub-unidade de Descritor de Tráfego - identifica o tipo de tráfego que será negociado, ou seja, o código do serviço (de acordo com o campo TOS do cabeçalho IP), portas de origem/destino e protocolo de rede utilizado;
- Sub-unidade de Descritor de Carga - caracteriza o condicionamento e policiamento do tráfego, através de parâmetros que reúnem tipo de descritor (mecanismo de policiamento), parâmetros relativos ao descritor (tamanho da rajada, vazão media, entre outros), e tratamento de excesso (condicionamento, remarcação e descarte).

- Sub-unidade de Parâmetros de QoS - identifica os parâmetros de QoS propriamente dito, ou seja, o atraso, a variação de atraso e o descarte máximo permitidos pelo contrato;
- Sub-unidade de Disponibilidade de serviços - identifica os horários de disponibilidade de QoS;
- Largura de Banda - identifica a largura de banda máxima que o cliente pode usufruir com a QoS prometida e delimitada pelas outras sub-unidades.

4.3 Arquitetura de Políticas

Através do mecanismo de arquitetura de Políticas é possível automatizar as tomadas de decisões do gerente de controle do BB em relação às diferentes condições da rede em que se deseje monitorar o tráfego, deixando a plataforma mais ágil e em melhores condições de atender aos requisitos de QoS impostos pelo cliente da rede.

Para esta dissertação foram propostas novas políticas de “Obrigações”, ver 2.7.1. Para o melhor manuseio das políticas, as mesmas foram divididas em três categorias: Inicialização, Atuação e de Negociação [Ped04].

4.3.1 Políticas de Inicialização

As políticas de Inicialização atuam nos roteadores de forma a definir:

- criação, seguida de configuração, prévia de LSP;
- número de filas físicas a serem inicializadas pelo roteador;
- qual o mecanismo de gerência de filas;
- tamanho máximo da filas;
- peso dos serviços já padronizados pela plataforma.

Para esta dissertação, de forma a atender os cenários de simulação, foram alterados os parâmetros de configuração do tamanho da fila e configuração de LSP. Este dados serão repassados no próximo capítulo 5 com detalhes quando serão relatados as simulações impostas à plataforma.

Vale ressaltar que as políticas de Inicialização são facilmente alteradas, uma vez que estão armazenados em arquivos do tipo texto.

4.3.2 Políticas de Atuação

As políticas de Atuação são aplicadas de maneira a corrigir uma quebra de contrato ou mesmo para prevenir uma violação de contrato de serviços. As políticas disponíveis na plataforma assumem o caráter de correção.

As novas políticas propostas têm o intuito de prevenir a violação dos contratos pré-estabelecidos, assim como corrigi-los.

As políticas que estavam disponíveis na plataforma monitoravam o nível de descarte de pacotes sofrido por um fluxo e o atraso para uma dada classe de serviço. Os tráfegos utilizados para verificação e validação eram de origem constante ou exponencial, não sendo utilizados tráfegos do tipo TCP [Gui01], [Ped04].

As novas políticas propostas vão levar em consideração dois aspectos do encaminhamento de pacotes, a largura de banda disponível e o nível de descarte. Não mais se levará em conta o atraso dos pacotes, pois como dito em 4.2.5, porém os mesmos dados serão levados em consideração durante a análise de dados.

Além dessas alterações, entre as ações a serem tomadas mediante quebra de contratos um tempo para estabilização será imposto ao sistema para que o mesmo possa se acomodar antes que uma outra ação seja tomada, evitando assim ações consecutivas para, praticamente, os mesmos dados de análise, levando a ações precipitadas de reconfiguração de filas ou desagregação de tráfego.

Para a política que visa o reestabelecimento de QoS para ao usuário final por meio da reconfiguração das filas dos roteadores, será também diminuído o número de tentativas de uma unidade antes que se atinja a ação de desagregar o tráfego, esta medida foi tomada após a introdução do tempo de espera pela estabilização do sistema.

4.3.3 Políticas Pré-existentes

No sistema estavam disponíveis duas políticas: uma primeira relacionada com o nível de descarte dos pacotes e uma segunda relacionada com o andamento dos níveis de atrasos dos pacotes.

- Política de Descarte

A política referente ao descarte de pacotes diz que: “Mediante um valor de descarte de pacotes maior que o tolerável estipulado por contrato para uma dada classe de serviço, o sistema deve reconfigurar as filas do roteadores de modo a fornecer uma banda maior para a referenciada COS. Se os limites de descartes se mantiverem acima do tolerável uma atualização no peso da banda é feito, segundo o algoritmo WRR, por mais três tentativas. E se após as quatro tentativas, o valor de descarte ainda não for tolerável, a classe de serviço que sofre descarte de pacotes é desviado para este novo LSP.”

Segue um pseudo-código da política citada acima:

Se o número de pacotes descartados por uma CoS maior do que o máximo permitido por contrato
Se for a primeira ocorrência (primeira tentativa)
 ação: Configura-se as filas físicas e virtuais com o peso de 10 unidades, seguindo o mecanismo WRR, para atender a classe de serviço que registrou perda de pacotes
Se número de ocorrência for menor do que 4 tentativas & maior que 1 tentativa
 ação: aumenta-se o peso na fila da classe de serviço em 10 unidades
Se número de tentativas for maior que 4
 ação: desagregação de tráfego para a classe de serviço, em questão, por meio da criação de um novo LSP.

- Política de Descarte

A política referente ao atraso de pacotes diz que: “Caso haja um atraso maior que o permitido por contrato, o sistema irá verificar se já houve alguma perda de pacotes e, por consequência, uma primeira reconfiguração na fila dos roteadores, em caso positivo, um novo LSP é criado para encaminhar o tráfego com atraso.”

Segue um pseudo-código da política citada acima:

Se for registrado atraso de pacotes para uma dada CoS
 for maior do que o máximo permitido por contrato & já houve no sistema uma reconfiguração nas filas do roteadores, ou seja, perda de pacotes acima do tolerável por contrato
 ação: desagregação de tráfego para a classe de serviço, em questão, por meio da criação de um novo LSP.

Nota: Esta política possui dependência direta da política de descarte.

4.3.4 Reformulação/Criação de Políticas de Atuação

As políticas de atuação já existentes sofrerão alterações, no sentido de não só controlar o nível de descarte dos pacotes das classes de serviços do LSP, mas também a ocupação de banda dos LSPs existentes e estipular um tempo para a estabilização do sistema para que ações sequenciais não atuem sobre, praticamente, os mesmos dados.

As novas políticas em questão também possuem o caráter de prevenção, assim como o de correção, pois, enquanto que as políticas de correção tem como objetivo observar

que o contrato de QoS foi quebrado e tentar recuperá-lo no menor tempo possível, as políticas de prevenção que tem como meta prever e impedir uma violação de contrato ou um agravamento da situação de desacordo com o usuário final.

Novas Políticas de Atuação

- Política de Correção

A política, a seguir, tem o caráter de correção. A política diz que: “Caso alguma classe de serviço tenha o seu patamar de descarte tolerável ultrapassado, a ação para reaver a QoS é rearranjar as filas dos roteadores por três tentativas, caso as três tentativas não sejam suficientes a classe de serviço que está sofrendo o descarte não desejado sofre uma desagregação, ou seja, ou o fluxo é re-roteado para outro LSP ou um novo LSP é criado para abrigar esta classe de serviço. A cada tentativa um ”*timer*“ é inicializado, somente ao expirar deste ”*timer*“ uma nova ação que seja imposta ao sistema pode ser aplicada.”

Segue um pseudo-código do conjunto de políticas citado acima:

Se o número de pacotes descartados por uma CoS for maior do que o máximo permitido por contrato

Se for a primeira ocorrência (primeira tentativa)

ação1: Configura-se as filas físicas e virtuais com o peso de 10 unidades, seguindo o mecanismo WRR, para atender a classe de serviço que registrou perda de pacotes

ação2: um ”*timer*“ é inicializado

Se número de ocorrência for menor do que 3 tentativas & maior que 1 tentativa & o ”*timer*“ expirou

ação1: aumenta-se o peso na fila da classe de serviço em 10 unidades

ação2: um ”*timer*“ é inicializado

Se número de tentativas for maior que 3 & o ”*timer*“ expirou

ação: procura-se um novo LSP para o fluxo com menor prioridade

Se encontrar um LSP adequado para este fluxo de dados

ação: agrega o tráfego a outro LSP

Senão

ação: desagrega o tráfego de menor prioridade através da criação de outro LSP

- Políticas de prevenção

Em paralelo com a política apontada acima, temos as políticas de monitoramento da ocupação de banda do LSP. As decisões tomadas por tais políticas levam em conta não só os valores de banda ocupada pelo LSP, mas também os dados contidos na tabela LFT.

A primeira política diz que: "No caso da alocação de banda atingir um nível de ocupação maior do que 95%, mesmo que não haja descarte, uma desagregação de tráfego é proporcionada ao tráfego de menor prioridade, ou seja, os tráfegos correntes no LSP pertencem a diferentes classes de serviço."

Segue um pseudo-código:

```

Se a ocupação de banda do LSP for maior do que 95%
    & as classes de serviço do LSP são distintas
    ação: procura-se um novo LSP para o fluxo com menor prioridade
    Se encontrar um LSP adequado para este fluxo de dados
        ação: agrega o tráfego a outro LSP
    Senão
        ação: desagrega o tráfego de menor prioridade através da criação de outro LSP

```

A segunda política diz que: "Mediante uma alocação de banda de LSP maior do que 85% e os fluxos associados ao LSP pertençam a mesma classe de serviço, uma desagregação de tráfego é aplicada, por meio da criação de um novo LSP ou pela agregação do tráfego indesejado em um outro LSP que atenda as necessidades do mesmo."

As políticas de correção e prevenção apresentadas acima não teriam validade neste cenário, pois se baseiam sua idéia de rearranjo da configuração das filas dos roteadores, levando somente o PHB em consideração.

Na política acima foi estipulado um valor de tolerância 10% maior, pois diante da chance da ocorrência da perda de pacotes, a reconfiguração das filas se apresenta como uma solução mais rápida e barata.

Segue um pseudo-código:

```

Se a ocupação de banda do LSP for maior que 85% do reservado
    Se todas as classes de serviço associadas ao LSP são iguais
        ação1: procura-se um novo LSP para agregar o tráfego
        ação2: seleciona-se um tráfego da tabela LFT
        Se encontrar um LSP adequado para este fluxo de dados
            ação: agrega o tráfego seleciona a outro LSP
    Senão

```

ação: desagrega o tráfego selecionado por meio da criação de um LSP dedicado

4.3.5 Políticas de Negociação

A decisão de como atender ao pedido dos clientes que desejam adentrar em um domínio acontece por meio das políticas de negociação que estabelecem o padrão de qualidade que o cliente deseja ter durante o seu trajeto na rede, sem comprometer os contratos de serviços que estão em vigor.

As políticas de negociação estão divididas em dois grupos: as políticas pertencente ao domínio solicitado e as políticas pertencentes ao cliente em si.

As políticas de admissão, restrição de banda e renegociação dizem respeito ao domínio solicitante, enquanto que as políticas de aceitação de contratos são operadas pelo cliente.

Políticas de Admissão

A política existente na plataforma, atualmente, restringi o acesso de clientes não cadastrados aos recursos do domínio, através de uma lista de clientes com permissão. Dessa forma, apresenta-se como uma política rápida e eficiente, pois evita uma desnecessária troca de mensagem.

Nova Política de Admissão

Para esta dissertação, fez-se necessário uma nova política de admissão que irá se preocupar com os fluxos já estabelecidos em um LSP.

A política diz que: "Caso o tráfego entrante tenha o mesmo destino e a mesma classe de serviço de um fluxo já corrente no LSP escolhido para carregá-lo, deve-se criar um LSP dedicado para o mesmo ou agregá-lo a um LSP pré-existente que possa atender ao destino requisitado, porém de rota diferenciada, com fluxos de diferentes classes de serviço."

Segue um pseudo-código:

Se o LSP, escolhido para o tráfego entrante, só possua fluxos com o mesmo PHB e destino do mesmo

ação: procura-se um novo LSP com o mesmo destino e diferente rota, cujo fluxos possuem classes de serviço diferenciadas

Se encontrar

ação: agrega-se tráfego ao LSP encontrado

Senão

ação: cria-se um LSP dedicado a esse fluxo de dados

Essa política é muito importante para uma rápida ação em caso de quebra de contrato, pois as políticas de Atuação existentes não sabem tratar com rapidez este caso de estudo, pois como as classes de serviço e o destino dos tráfegos entrantes são iguais, as políticas que tratam o monitoramento da ocupação de banda de um LSP não seriam acionadas, somente as políticas de descarte de pacotes seriam acionadas, provendo resultados satisfatórios somente após o tempo de três tentativas de reconfiguração de filas.

Políticas de Restrição de Banda

De forma a não limitar os recursos de um domínio a um único cliente, faz-se necessária as políticas de restrição de banda.

As políticas de restrição de banda definem qual será a quantidade de recurso máxima que pode ser cedida para cada serviço e para cada domínio adjacente, seguindo o modelo Cascata.

Políticas de Renegociação

Com o intuito de controlar a negociação dos contratos de atacado, faz-se valer das políticas de renegociação que definem o tempo mínimo para a renegociação de um contrato em que se permite novas negociações após o tempo de validade do contrato ficar expirado.

Esta categoria de política de negociação não se faz aplicada para contratos do tipo varejo.

Políticas de Aceitação de Contratos

Este conjunto de políticas define como o cliente deve se comportar no caso de contra-propostas, podendo:

- não aceitar contra-proposta;
- aceitar qual seja a contra-proposta;
- aceitar uma quantidade mínima de recursos.

Capítulo 5

Experimentos/Validações

Este capítulo expõe os experimentos realizados na plataforma apresentada no Capítulo 4, cujas finalidades são:

1. verificar e validar uma nova tabela para a composição do mecanismo de “Engenharia de Tráfego” da plataforma;
2. verificar e validar novas políticas que fazem com que a plataforma, em questão, se torne mais dinâmica e eficiente no tratamento de situações de sobrecarga de recursos do sistema.

A topologia utilizada para os experimentos é visualizada pela Figura 5.1. Os hexágonos representam os clientes, enquanto que os círculos representam os roteadores MPLS/Diffserv que são configurados/gerenciados por um *Bandwidth Broker*(BB). Todo enlace da topologia possui uma capacidade de 3.456Mbps.

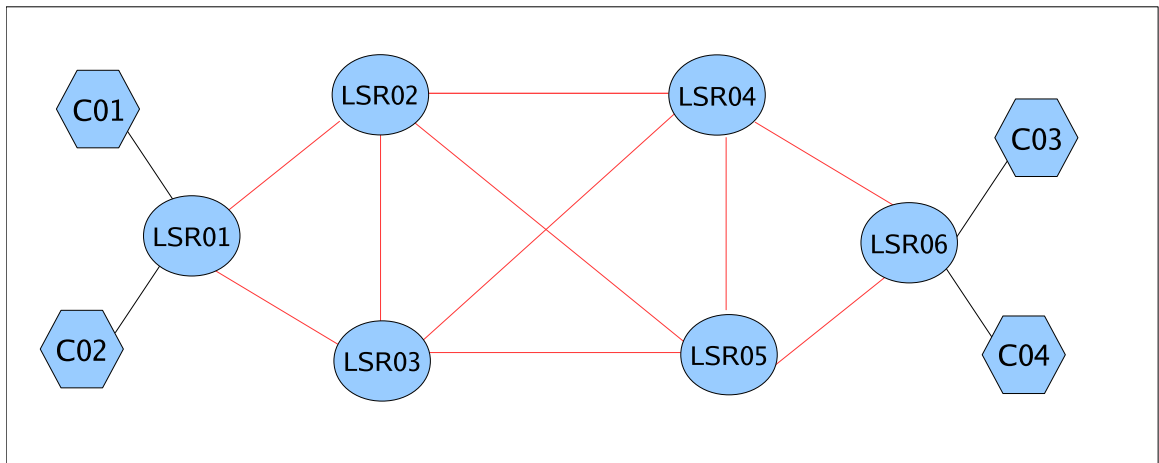


Figura 5.1: Topologia de rede usada nos experimentos

Para os experimentos será feito um monitoramento do tráfego que seguirá as condições expostas na sub-seção 4.2.5 e terá uma frequência de 0.5 segundo para acontecer.

Foram selecionados dois tipos de tráfego a serem utilizados nas simulações:

Tipo de Trafego	Tamanho do Pacote	COS	Código da COS
Multimídia	210bytes	EF	20
Voz	66bytes	EF	20
TCP	1000bytes	AF	30

Tabela 5.1: Características dos tráfegos gerados

Com relação às classes de serviço, foram utilizadas duas classes de serviço: *Assured Forwarding* (AF) e *Expedited Forwarding* (EF). Os dados relativos a atrasos máximos e descarte máximos permitidos para cada classe podem ser verificados na Tabela 5.2 [Gui01].

Classe	Atraso Máximo	Descarte Máximo
<i>Assured Forwarding</i>	4s	20%
<i>Expedited Forwarding</i>	210ms	20%

Tabela 5.2: Características das classes de serviço

As filas que estão associadas aos roteadores são do tipo RED com especialização em DiffServ. Dessa maneira, a estrutura de fila consiste de 4 filas físicas, cada qual contendo 2 filas virtuais. Cada fila física representa uma classe de serviço, onde cada associação de fila física e virtual está associada a um *codepoint*.

O tamanho das filas estipulado para os experimentos pré-existentes era de 1000 pacotes, sendo ajustado para 100 pacotes, para que fosse possível caracterizar os atrasos provocados pelas novas políticas de atuação.

Além dessa modificação, no momento de criação do LSP foi especificado dois destinos ao mesmo, ou seja, cliente 04 por meio do cliente 02 e cliente 03 por meio do cliente 01.

5.1 Experimento 1

O objetivo deste experimento é provar a eficiência da tabela LFT junto ao mecanismo de "Engenharia de Tráfego" proposto pela plataforma pré-existente com base no MNS [AC01].

Inicialmente na simulação é a inicialização o BB cria um E-LSP de identificação de número 2000 que pode encaminhar tráfego do cliente 01 ou cliente 02 para os clientes 03 e 04. A rota estipulada para este LSP é LSR01-LSR02-LSR04-LSR06.

Durante a inicialização do LSP é configurado que, pelo menos, será possível encaminhar os tráfego com classes de serviço EF e AF para os dois destinos possíveis do LSP.

Para encerrar a fase inicial de inicialização da plataforma, o monitoramento do tráfego é iniciado.

Após a fase de inicialização, chega a vez da negociação. No tempo de 1s, o cliente 01 negocia satisfatoriamente com o BB um contrato de 2Mbps para fluxos de tráfego da classe EF e 780Kbps para fluxos da classe de serviço AF até o cliente 03, logo, em seguida, o cliente 02 negocia os mesmos valores de fluxos de tráfego, com destino ao cliente 04.

No tempo de 5s, inicia-se um tráfego de vídeo do tipo CBR que pertence à classe de serviço EF com uma taxa de 1Mbps que se estende do cliente 01 ao cliente 03. Este tráfego se acomoda no LSP2000, como esperado, de uma forma harmoniosa, ou seja, não se verifica perda ou atraso dos pacotes. Este tráfego é registrado na tabela LFT, como mostra a Tabela 5.3.

LSP	Classe de Serviço	Destino
2000	20	14

Tabela 5.3: Tabela LFT

Já no tempo de 15s, inicia-se um novo tráfego de vídeo com uma taxa de 2Mbps no LSP 2000 cuja a classe de serviço também é do tipo EF que se estende do cliente 02 ao cliente 04. Este tráfego também é, devidamente, registrado na nova Tabela 5.4.

LSP	Classe de Serviço	Destino
2000	20	Cliente 03
2000	20	Cliente 04

Tabela 5.4: Tabela LFT

Durante o monitoramento deste tráfego, percebe-se que em $t=15,1s$ a ocupação de banda do LSP 2000 ultrapassa a porcentagem dos 85%, este patamar faz com que a tabela LFT seja consultada e uma vez verificada que a política de correção baseada em filas não irá ter serventia para este cenário, pois todos os tráfegos são do tipo EF, a política de prevenção descrita em 4.3.4 é disparada.

Dessa maneira, um novo LSP de identificação 2999 e rota R0-R2-R4-R6-R7 é criada e o tráfego cujo o destino é o cliente 03 é desviado para este LSP. A criação de um novo LSP é obrigatória neste caso, pois não há um outro LSP criado que atenda as necessidades do tráfego desviado. Este dinamismo pode ser acompanhado pelos gráficos expostos na Figura 5.2.

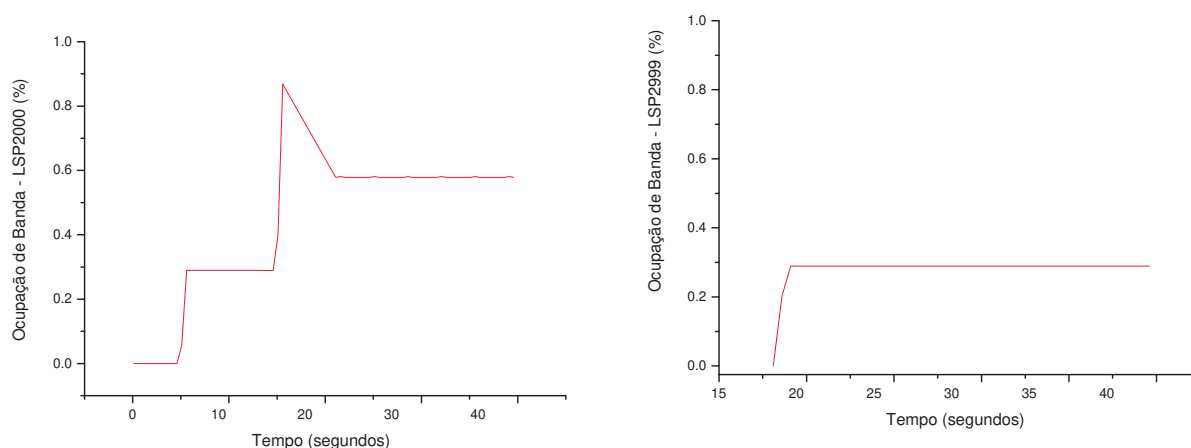


Figura 5.2: Porcentagem de ocupação de banda para os LSPs 2000 e 2999

No tempo 19,1s, verifica-se que o tráfego do cliente 01 para o cliente 03 passa a ser roteado para o LSP 2999, ou seja, o LSP200 deixa de ficar super alocado e o LSP2999 passa a receber todo o tráfego desviado. E em 23,6s, o LSP 2000 conta só com o tráfego do cliente 02 ao cliente 04.

Dessa maneira, não se verifica descarte durante o decorrer da simulação, assim como alocação excessiva de recursos do sistema. E os atrasos nas filas se mantêm bem abaixo dos níveis de tolerância estipulados pelos contratos, como mostra a Figura 5.3.

Para o tráfego do cliente 01 ao cliente 03, verifica-se um atraso pequeno da ordem de 24ms até o tempo de 15 segundos, ou seja, até o momento de início do tráfego do cliente 02 ao cliente 04, tendo um pico em 15,5s de aproximadamente 28ms e voltando à estabilização em 16,5, caracterizando um período de pico de 1 segundo. Enquanto que o tráfego do cliente 02 ao cliente 04 entra em regime em 15,5s, com um atraso de aproximadamente 31ms.

Através deste experimento, pode-se concluir que em um período de 4 segundos foi possível a detecção e aplicação de uma medida de prevenção no sistema, evitando a possibilidade de quebra de contrato. O que através das políticas de correção só seria

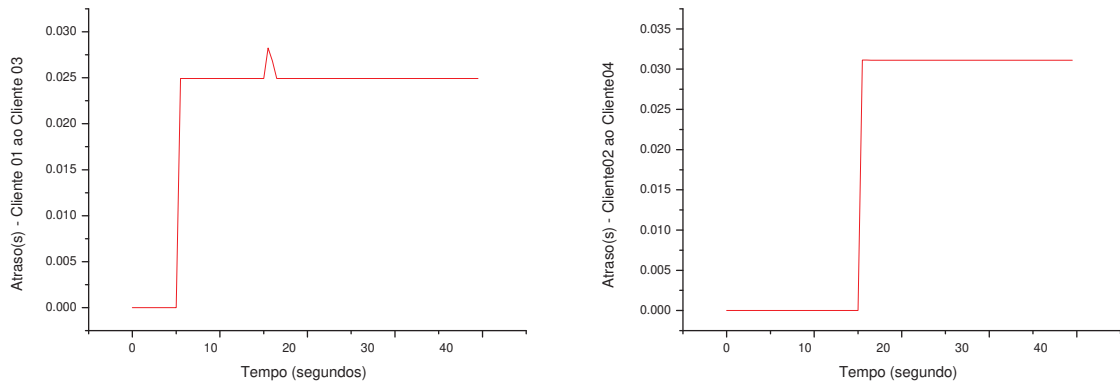


Figura 5.3: Medidas de atraso dos cliente 01 ao Cliente 03 e do cliente 02 ao cliente 04

detectado quando a ocupação de banda estivesse na ordem de 95% ou quando o sistema estivesse já operando com perda de pacotes.

5.2 Experimento 2

Este experimento tem por objetivo mostrar a otimização feita na política de atuação já existente na plataforma. A nova política leva em consideração não só os níveis de descarte dos pacotes para uma dada classe de serviço, mas também a ocupação de banda do LSP, chegando a uma tomada de decisão mais eficiente, em um tempo menor.

O experimento tem início com a criação de um ELSP de identificação de número 2000, por parte do BB, que possui como origem os clientes 01 e cliente 02 e destinos os clientes 03 e 04. A rota estipulada para este LSP é LSR01-LSR02-LSR04-LSR06.

Durante a inicialização do LSP, configura-se que pelo mesmo será possível encaminhar os tráfegos com classes de serviço EF e AF para os dois destinos possíveis do LSP.

Para encerrar a fase de inicialização da plataforma, o monitoramento do tráfego é iniciado.

Após a fase de inicialização, chega a vez da negociação. No tempo de 1s, o cliente 01 negocia satisfatoriamente com o BB um contrato de 1.5Mbps para fluxos de tráfego da classe EF e 2.780Mbps para fluxos da classe de serviço AF até o cliente 03, logo, em seguida, o cliente 02 negocia os mesmos valores de fluxos de tráfego, com destino ao cliente 04.

No tempo de 5s, inicia-se um tráfego FTP que pertence à classe de serviço AF que se estende do cliente 02 ao cliente 04. O tráfego gerado é registrado na tabela LFT, vide

Tabela 5.5.

LSP	Classe de Serviço	Destino
2000	30	15

Tabela 5.5: Tabela LFT - tráfego FTP

No tempo de 15s, inicia-se um novo tráfego de voz com uma taxa de 1Mbps no LSP2000 cuja classe de serviço é EF que se estende do cliente01 ao cliente 03. Este tráfego também é, devidamente, registrado na Tabela 5.6.

LSP	Classe de Serviço	Destino
2000	30	15
2000	20	14

Tabela 5.6: Tabela LFT - tráfegos FTP e Voz

Em 15,1s, o monitor de tráfego emite um alarme indicando que o nível de descarte para o tráfego EF ultrapassou o limite máximo estipulado pelo contrato. Neste momento a tabela LFT é consultada e se constata que a política de configuração de filas pode vir a estabilizar o sistema sem haver a necessidade da criação de um novo LSP, uma vez que os tráfegos correntes no LSP são de diferentes classes de serviço.

Após esta dedução, o sistema aciona a política de atuação baseada em configuração de filas, onde a primeira atitude é configurar as filas físicas e virtuais dos roteadores de forma a diferenciar os tráfegos AF e EF. O tráfego EF será tratado pela fila física 1 e virtual 0 com um peso de 10 unidades, segundo o algoritmo WRR.

Além disso, o tráfego TCP se retraiu a ponto de parar de transmitir, como mostra Figura 5.4, no intervalo de 16,5s a 17,5 segundo.

No período de retração do tráfego TCP, acontece perda de pacotes que ultrapassa o limite tolerável estabelecido pelo contrato, porém nenhuma ação é tomada, pois entre uma ação de reconfiguração de fila e outra existe um tempo de 3 segundos para que o sistema possa tender à estabilização. A partir do momento 17,5s, o tráfego TCP recomeça a ser transmitido e os níveis de descarte voltam a ter valores aceitáveis, vide Figura 5.5. Pela mesma figura também é possível acompanhar a curva de pacotes duplicados que acompanha o nível de descarte de pacotes.

No tempo de 19,6s, o nível de descarte de pacotes do PHB 20 é novamente ultrapassado e uma segunda reconfiguração de filas é realizada de forma que o peso da classe de serviço EF é acrescido em 10 unidades, segundo o algoritmo WRR.

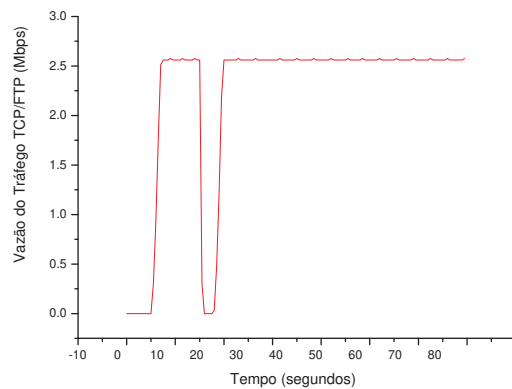


Figura 5.4: Vazão do tráfego FTP

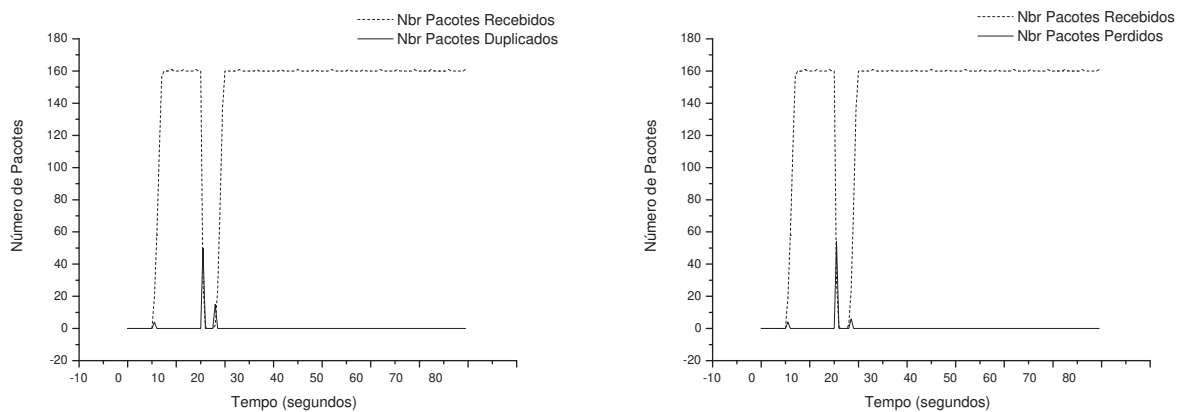


Figura 5.5: Número de Pacotes Recebidos, Duplicados e Perdidos

Porém, esta medida continua a não apresentar bons resultados, pois, logo em seguida a esta ação, os níveis de descarte não cedem como mostra a Figura 5.6.

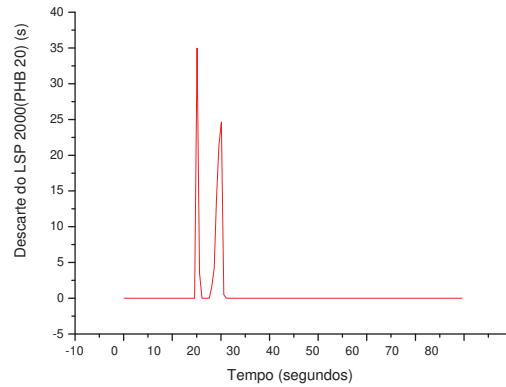


Figura 5.6: Descarte de pacotes para o LSP2000 - classe de serviço EF

Porém, no tempo de 20.1s, verifica-se uma alocação de banda do LSP 2000 acima de 95%. Para que a banda não seja estourada, a política de criação de um novo LSP é acionada, sendo assim, um novo LSP é criado imediatamente. Este novo LSP passa a atender a demanda de fluxos de menor prioridade, segundo a tabela LFT. A política entende que deve deslocar o tráfego AF para que as características de contratos impostas ao início da simulação não deixem de ser atendidas.

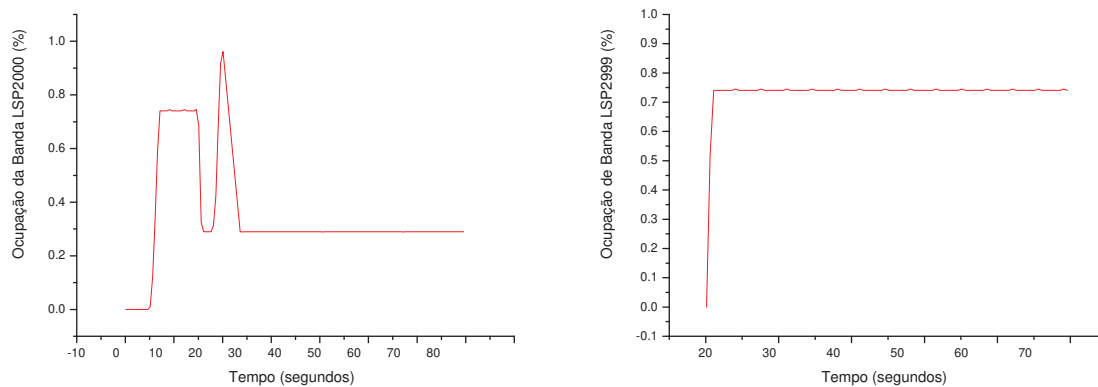


Figura 5.7: Níveis de ocupação de banda para os LSPs 2000 e 2999

Dessa maneira, o LSP 2999 é criado com rota LSR01-LSR02-LSR04-LSR06, desafogando o LSP 2000, através do encaminhamento do tráfego do cliente 02 ao cliente 04, como mostra a Figura 5.7.

Assim como a vazão do tráfego FTP volta à sua taxa de 320 pacotes por minutos, sem perda de pacotes ou registro de pacotes duplicados, vide gráficos da Figura 5.5. Os gráficos mostram valores de 160 pacotes por 0.5 s.

Ao final da simulação, verifica-se que a criação de um LSP para desafogar o tráfego com descarte de pacotes quando o sistema está prestes a sobrecarregar um recurso é uma boa solução mediante soluções mais brandas como a reconfiguração do peso das classes de serviços nas filas dos roteadores.

5.3 Experimento 3

O experimento 3 visa verificar o tempo de estabilização requisitado entre as ações de reconfiguração dos pesos das filas, mediante um cenário que não há a ultrapassagem do limite máximo possível de ocupação de banda do LSP, mas descarte de pacotes para a classe de serviço EF.

Assim como nas simulações anteriores, no momento inicial do experimento, um LSP do tipo ELSP é criado pelo BB com a identificação 2000, possuindo origem no LSR01 e destino os C03(cliente 3) e C04(cliente 4).

A rota estipulada para este LSP é LSR01-LSR02-LSR04-LSR06. O cliente 03 será mapeado pelo sistema com o número 14 e o cliente 04 pela identificação de número 15.

Ainda durante a inicialização do LSP, configura-se no mesmo que será possível o encaminhamento dos tráfegos com classes de serviços EF e AF para os dois destinos possíveis do LSP.

Para encerrar a fase de inicialização da simulação, o monitoramento do tráfego é iniciado.

Após a fase de inicialização, chega a vez da negociação. No tempo de 1s, o BB negocia, com sucesso, fluxos de tráfego AF de 1.456Mbps e EF de 2Mbps para a classe EF.

No tempo de 5s, inicia-se um tráfego FTP que pertence à classe de serviço AF que se estende do cliente 02 ao cliente 04, sendo devidamente registrado na Tabela 5.7. Até o momento de 15s não é verificado atraso, descarte ou uso excessivo dos recursos de rede.

LSPs	PHB	DEST
2000	30	15

Tabela 5.7: Tabela LFT - Adição do primeiro fluxo da simulação (experimento 3)

A partir do tempo de 15s é iniciado um tráfego de vídeo com a taxa constante de 768kbps e classe de serviço EF, sendo registrado na tabela LFT, vide Tabela 5.8.

LSPid	PHB	DEST
2000	30	15
2000	20	14

Tabela 5.8: Tabela LFT - Adição do segundo fluxo da simulação (experimento3)

Já em 15,1s o sistema reporta um excesso de descarte de pacotes para a dada classe de serviço da ordem de aproximadamente 34,1%. Como esperado, a tabela LFT foi consultada e como ambos os tráfegos possuíam classes de serviço distintas, a ação a ser tomada foi a reconfiguração de fila para a classe EF.

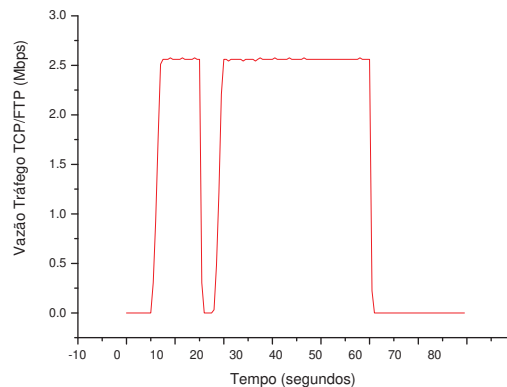


Figura 5.8: Vazão do tráfego FTP

Após essa primeira reconfiguração, o sistema apresentou uma certa melhora, porém não estável, pois no tempo de 19,6s houve uma nova perda excessiva de pacotes. Vale ressaltar que a melhora no nível de descarte do sistema também ocorreu, pois o tráfego TCP/FTP se retraiu totalmente com a introdução do tráfego CBR. Este fenômeno é bem observado pelo gráfico de vazão do tráfego FTP da Figura 5.8, onde no tempo 16.1s até o tempo 17.6s só houve a presença do tráfego de voz no LSP.

No período de tempo de 16.1s até 17.6s, o tráfego FTP aponta a ocorrência de pacotes duplicados e perdidos, vide gráficos exibidos na Figura 5.9.

Com a reestruturação do tráfego TCP/FTP, em 19.6s, um novo relato de excesso de perda de pacotes da ordem de aproximadamente 22,28% é verificado para a classe

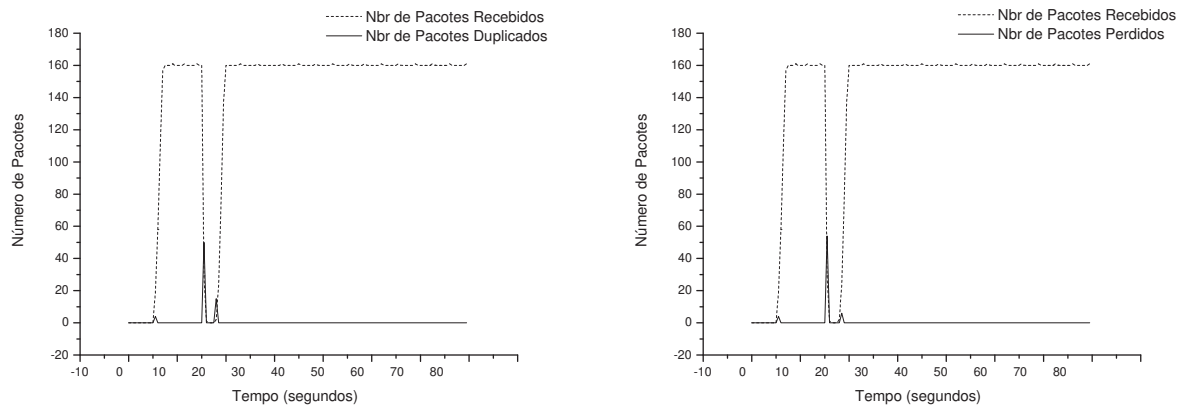


Figura 5.9: Vazão do tráfego FTP

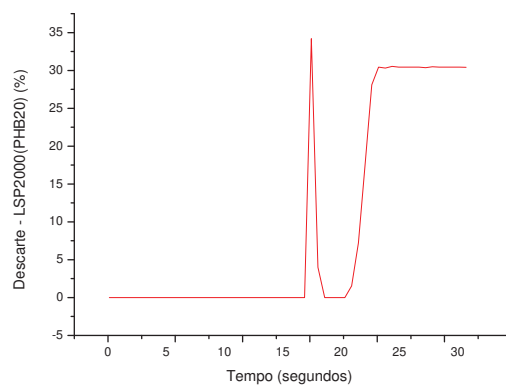


Figura 5.10: Descarte de pacotes para o LSP2000 - classe de serviço EF

de serviço EF. Para este caso, uma segunda tentativa de reconfiguração de fila é feito, aumentando a porcentagem de banda para o tráfego EF, porém não se verifica uma melhora nas taxas de perda de pacote, com isso em 23,1s uma última tentativa é aplicada e, novamente, não se constata uma redução nos níveis de perda de pacotes, como mostra gráfico da Figura 5.10.

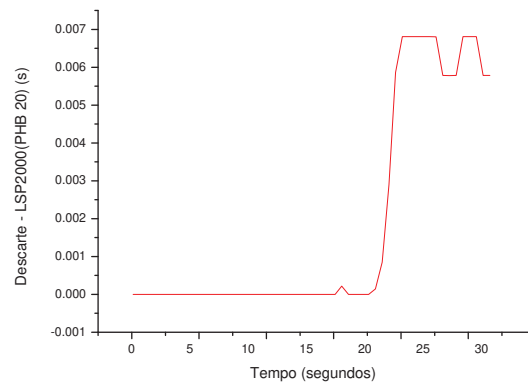


Figura 5.11: Atraso de pacotes para o LSP2000 - classe de serviço EF

As taxas de atraso para o PHB acompanharam as taxas de atraso, porém sem ultrapassar as taxas toleráveis pré-definidas por contratos, vide Figura 5.11.

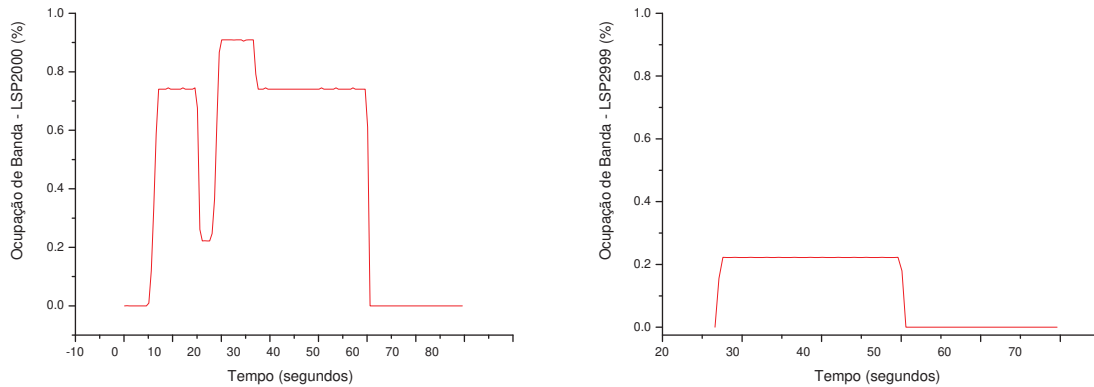


Figura 5.12: Níveis de ocupação de banda para os LSPs 2000 e 2999

Depois das três tentativas expostas acima, com período de espera de estabilização da

ordem de 3s, chega-se à conclusão que a melhor alternativa é a criação de um novo LSP para a classe de serviço EF, o que acontece no tempo 26.6 segundos de experimento. O novo LSP tem a identificação de LSP2999 e rota R0-R2-R4-R6-R7.

Durante o período de simulação não houve uma ocupação de banda do LSP maior que 95%, como mostram os gráficos da Figura 5.12.

Com este experimento, verificou-se que as três tentativas de reconfiguração de peso de filas, em um período um pouco maior do que 7 segundos, foram suficiente para constatar que para tráfegos intensos, a desagregação de tráfego é uma boa solução.

5.4 Considerações Finais sobre as Simulações Realizadas

Neste capítulo foram apresentadas três simulações, através das quais, foi feito um estudo do comportamento de tráfegos amplamente utilizados nas redes atuais em uma plataforma de entendimento complexo, porém de fácil configuração.

Os experimentos tinham como objetivo maior validar as novas políticas de gerenciamento de rede. As novas políticas contaram com o apoio de uma nova tabela LFT para o módulo MPLS-DS do NS, assim como um *timer* para estabilizar o sistema antes da próxima ação e as medidas de ocupação de banda como acionadoras de políticas.

Os cenários das simulações apresentados tinham a intenção de sobrecarga do sistema, assim como quebra de contratos pré-estabelecidos durante os momentos iniciais das simulações.

Os resultados obtidos foram válidos e abrem caminho para uma exploração ainda maior da plataforma, cuja bases de comunicação entre os elementos remotos está estável.

Aspectos de renegociação não foram explorados nas políticas e experimentos expostos.

Capítulo 6

Conclusão e Trabalhos Futuros

6.1 Conclusão

Este trabalho contribui com o aperfeiçoamento e criação de políticas de gerenciamento de rede, bem como com a criação de uma nova estrutura de "Engenharia de Tráfego". Essas melhorias foram aplicadas a uma plataforma MPLS com suporte à diferenciação de serviços.

Os resultados obtidos para os estudos de casos com tipos de tráfegos amplamente encontrados na Internet, em situações de sobrecarga, se mostraram satisfatórios, ou seja, os contratos estabelecidos com os clientes na admissão do tráfego foram respeitados, garantindo dessa maneira QoS de fim-a-fim. Ao garantir QoS aos clientes participantes das simulações, a plataforma validou as propostas de políticas e, conseqüentemente, um novo mecanismo de "Engenharia de Tráfego", que conta com a introdução de mais uma tabela para gerenciar as informações relacionadas ao LSP. Esta nova tabela foi denominada LFT, pela qual é possível diferenciar os fluxos de um LSP não só pelo PHB, mas também pelo destino dos pacotes.

Diversos trabalhos foram realizados com o intuito de prover políticas de prevenção ao sistema, no tocante ao monitoramento da ocupação de banda de um LSP e outros trabalhos se disponibilizam a oferecer políticas de correção levando em consideração o acompanhamento do nível de descarte e atraso dos pacotes. Outro aspecto levado em consideração nesta dissertação foi o fato de propor novas políticas de atuação com caráter tanto preventivo como de correção, nas quais não só o descarte dos pacotes, mas também o nível de ocupação da banda foram monitorados, oferecendo à plataforma um grau maior de abstração e dinamismo.

A introdução da nova estrutura de "Engenharia de Tráfego" permitiu um acompanhamento do tráfego com maiores detalhes, possibilitando a implementação de políticas de gerenciamento de tráfego mais sofisticadas. Esta nova estrutura pode ser encarada

como complementar ao módulo MPLS-DS do NS-2 dedicado a redes MPLS com suporte a serviços diferenciados, podendo ser utilizada em pesquisas destinadas ao assunto.

Através da exposição da plataforma a novos estudos de casos, pode-se concluir que a combinação dos modelos DiffServ e MPLS apresenta boas perspectivas no tocante à obtenção de QoS fim-a-fim para redes IP. No particular, a plataforma, embora complexa, se mostrou flexível às novas implementações e configurações.

A plataforma também se mostrou disponível para sua utilização por empresas que desejem utilizar um ambiente de simulação de novas políticas administrativas. Dessa maneira, ao estabelecer um contrato com um novo cliente a empresa que possui uma ferramenta de simulação pode se preparar com antecedência para certas situações particulares de tráfego referente ao novo cliente, assim como para responder às possíveis combinações do novo perfil de tráfego com os perfis já existentes e testados na rede.

6.2 Trabalhos Futuros

No tocante à plataforma, verifica-se a necessidade de adaptá-la a um modelo mais portátil, ou seja, menos dependente da plataforma Unix, apesar da dependência direta do módulo NS ao sistema operacional. Uma proposta para esta questão seria a implementação do *Bandwidth Broker* em Java, dessa maneira, a atualização da plataforma base ficaria a cargo das atualizações freqüentes apenas do módulo NS.

Outro ponto a melhorar na plataforma é a questão da modelagem e armazenamento de políticas. Uma possível solução para este ponto é o uso de servidores LDAP, o que possibilitaria um crescimento do número de contratos e políticas, sem que os mesmos se tornem um gargalo para o sistema.

Com relação às políticas, verifica-se a necessidade de ampliar as políticas de renegociação, de maneira a permitir a renegociação dos parâmetros de qualidade que o cliente deseje não só na admissão do tráfego, porém durante o seu trajeto na rede. Essa renegociação poderia se dar de maneira automática, ou seja, o tráfego ao adentrar em um domínio visto que o mesmo possua recursos em abundância para atender o mesmo, o mesmo pode ser configurado com maiores privilégios, ao passo que à medida que os recursos se tornem mais escassos, os critérios de qualidade podem voltar a valores mínimos de atendimento.

Um segundo aspecto a ser explorado na plataforma são estudos de casos que englobem o gerenciamento de filas que utilizam o algoritmo RED e suas variações. A abordagem poderia ser no sentido de possibilitar uma armazenagem do histórico das filas para que uma configuração mais precisa por parte dos administradores de rede fosse possível em diferentes ocasiões.

Outro aspecto que merece atenção em trabalhos futuros junto a essa plataforma é o de

políticas que sejam capazes de gerenciar conflitos de políticas, principalmente, na questão da desagregação de tráfego. O sistema pode operar de maneira impeditiva no caso de duas políticas desejarem desagregar tráfegos, ou seja, antes do sistema executar a ação de desagregação de tráfego, o mesmo deve verificar se já existe uma desagregação em andamento e por quanto tempo. Após esperar o fim com sucesso da desagregação por um período para a estabilização do sistema, uma nova checagem dever ser feita para saber se o pedido de desagregação, anteriormente disparado, continua válido.

Referências Bibliográficas

- [AC01] Gaeil Ahn and Woojik Chun. Design and implementation of mpls network simulator (mns) supporting qos. *15th International Conference on Information Networking, 2001*, pages 694–699, 2001.
- [AMA⁺99] D. Awduche, J. Malcolm, J. Agogbua, M. O'Dell, and J. McManus. Requirements for traffic engineering over mpls. (RFC 2702), Setembro de 1999.
- [Arc01] Bandwidth Broker Architecture. Qbone bandwidth broker architecture. Technical report, 2001. <http://qos.internet2.edu/qos/wg/documents-informational/20020709-chimne%to-et-al-qbone-signaling/>. Acessado em Outubro de 2006.
- [ASO05] Tricha Anjali, Caterina Scoglio, and Jaudelice Cavalcante Oliveira. New mpls network management techniques based on adaptive learning. *IEEE Transactions on Neural Networks*, 16(5):1242–1255, 2005.
- [BCB06] J. Babiarz, K. Chan, and F. Baker. Configuration guidelines for diffserv service classes. Rfc 4594, Agosto 2006.
- [BQ01] M. Brunner and J. Quittek. Mpls management using policies. *International Symposium on Integrated Network Management Proceedings, 2001 IEEE/I-FIP*, pages 515–528, 2001.
- [CHL⁺03] Ji-Feng Chiu, Zuo-Po Huang, Chi-Wen Lo, Wen-Shyang Hwang, and Ce-Kuen Shieh. Supporting end-to-end qos in diffserv/mpls networks. *10th International Conference on Telecommunications*, 1:261–266, Março de 2003.
- [CSD⁺01] K. Chan, J. Seligson, D. Durham, S. Gai, K. McCloghrie, S. Herzog, F. Reichmeyer, R. Yavatkar, and A. Smith. Cops usage for policy provisioning (cops-pr). Rfc 3084, Março 2001.

- [DCB⁺02] B. Davie, A. Charny, J.C.R. Bennet, K. Benson, J.Y. Le Boudec, W. Courtney, S. Davari, V. Firoiu, and D. Stiliadis. An expedited forwarding phb (per-hop behavior). Rfc 3246, Março 2002.
- [dOSAU04] Jaudelice C. de Oliveira, Caterina Scoglio, Ian F. Akyildiz, and George Uhl. New preemption policies for diffserv-aware traffic engineering to minimize re-routing in mpls networks. *IEEE/ACM Transactions on Networks*, 12(4):733–745, Agosto de 2004.
- [FWD⁺03] F. Le Faucheur, L. Wu, B. Davie, S. Davari, P. Vaananen, R. Krishnan, P. Cheval, and J. Heinanen. Multi-protocol label switching (mpls) support of differentiated services. Rfc 3270, IETF, Maio 2003.
- [Gro] Object Management Group. <http://www.omg.org>. Acessado em janeiro de 2007.
- [Gro02] D. Grossman. New terminology and clarifications for diffserv. Rfc 3260, Abril 2002.
- [Gui01] Alex Marcelo Samaniego Guillén. Mpls-ds: Uma plataforma para validação de políticas no contexto das redes mpls/diffserv. Master’s thesis, FEE, Unicamp, Dezembro 2001.
- [HBWW99] J. Heinanen, F. Baker, W. Weiss, and J. Wroclawski. Assured forwarding phb group. Rfc 2597, Junho 1999.
- [HPW02] D. Harrington, R. Presuhn, and B. Wijnen. An architecture for describing simple network management protocol (snmp) managemen frameworks. Rfc 3411, Dezembro 2002.
- [ITC] ITCL. tcltk.com. <http://www.tcltk.com/itcl>. Acessado em janeiro de 2007.
- [KLHm04] Yu Ke, Zhang Lin, and Zhang Hui-min. A preemption-aware path selection algorithm for diffserv/mpls networks. *Proceedings IEEE Workshop on IP Operations and Management, 2004*, pages 129–133, Outubro de 2004.
- [MC01] Maurício Ferreira Magalhães and Eleri Cardozo. *Introdução à Comutação IP por Rótulos Através de MPLS*. SBRC de 2001, Maio de 2001.
- [Mon] Sink Monitor. Luigis’s homepage. Technical report, Lab. d’Informatique Paris 6 - Networks and Performance Analysis. <http://www-rp.lip6.fr/~iannone/>. Acessado em janeiro de 2007.

- [Moy98] J. Moy. Ospf version 2. Rfc 2328, Abril 1998.
- [MSZ96] W. Ma, P. Steenkiste, and H. Zhang. Routing high-bandwidth traffic in max-min fair share network. *Proceedings of the ACM SIGCOMM*, pages 206–217, Agosto de 1996.
- [NC01] K. Nichols and B. Carpenter. Definition of differentiated services per domain behaviors and rules for their specification. Rfc 3086, Abril 2001.
- [NS01] NS. The network simulator - ns-2, 2001. <http://www.isi.edu/nsnam/ns>. Acessado em janeiro de 2007.
- [Ped04] Rodrigo Neiva Pedatella. Um modelo de negociação de serviços com qos fim-a-fim na internet. Master's thesis, IC, Unicamp, 2004.
- [Pie00] P. Piedad. A network simulator differentiated services implementation. 2000.
- [RVC01] E. Rosen, A. Viswanathan, and R. Callon. Multiprotocol label switching architecture. Rfc 3031, Janeiro 2001.
- [RVKF99] Raju Rajan, Dinesh Verma, Sanjay Kamat, and Eyal Felstaine. A policy framework for integrated and differentiated services in the internet. *IEEE Network*, 13(5):36–41, Setembro/Outubro de 1999.
- [SL02] M. Solman and E. Lupu. Security and management policy specification. *IEEE Network*, 16(2):10–19, Março de 2002.
- [SSJ02] B. Szviatovsky, A. Szentezi, and A. Juttner. Minimizing re-routing in mpls networks with preemption-aware constraint-based routing. *Computer Communications*, (25):1076–1083, 2002.
- [Tan04] Andrew S. Tanenbaum. *Computer Network - Fourth Edition*. 2004.
- [TCL] CORBA Scripting With TCL. Combat. <http://www.fpx.de/Combat/>. Acessado em janeiro de 2007.
- [WHK06] M. Wahl, T. Howes, and S. Kille. Lightweight directory access protocol (ldap) : Directory information models. Rfc 2702, Setembro 2006.