

Gerência de Redes Virtuais Privadas de Camada 1

Este exemplar corresponde à redação final da Dissertação devidamente corrigida e defendida por Neumar Costa Malheiros e aprovada pela Banca Examinadora.

Campinas, 29 de agosto de 2006.

Edmundo Roberto Mauro Madeira
(Orientador)

Dissertação apresentada ao Instituto de Computação, UNICAMP, como requisito parcial para a obtenção do título de Mestre em Ciência da Computação.

UNIDADE BC
 Nº CHAMADA:
TI/UNICAMP
M294g.
 V. _____ Ed. _____
 TOMBO BC/ 72735
 PROC. 16.115-07
 C ☐ D ☒
 PREÇO 11,00
 DATA 30/05/07
 BIB-ID 42527

**FICHA CATALOGRÁFICA ELABORADA PELA
 BIBLIOTECA DO IMECC DA UNICAMP**

Bibliotecária: Miriam Cristina Alves – CRB8a / 5094

Malheiros, Neumar Costa

M294g Gerência de redes virtuais privadas de camada 1 / Neumar Costa
 Malheiros -- Campinas, [S.P. :s.n.], 2006.

Orientador : Edmundo Roberto Mauro Madeira

Dissertação (mestrado) Universidade Estadual de Campinas,
 Instituto de Computação.

1. Redes de computadores (Gerenciamento). 2. Redes de
 computação – Arquitetura. 3. Redes de computação - Protocolos. I.
 Madeira, Edmundo Roberto Mauro. II. Universidade Estadual de
 Campinas. Instituto de Computação. III. Título.

Título em inglês: Layer 1 virtual private network management

Palavras-chave em inglês (Keywords): 1. Computer network (Management). 2. Computer
 network - Architecture. 3. Computer network - Protocol.

Área de concentração: Sistemas de Computação

Titulação: Mestre em Ciência da Computação

Banca examinadora: Prof. Dr. Edmundo Roberto Mauro Madeira (IC/UNICAMP)
 Prof. Dr. Djamel Fawzi Hadj Sadok (CIn/UFPE)
 Prof. Dr. Nelson Luis Saldanha da Fonseca (IC/UNICAMP)

Data da defesa: 29/08/2006

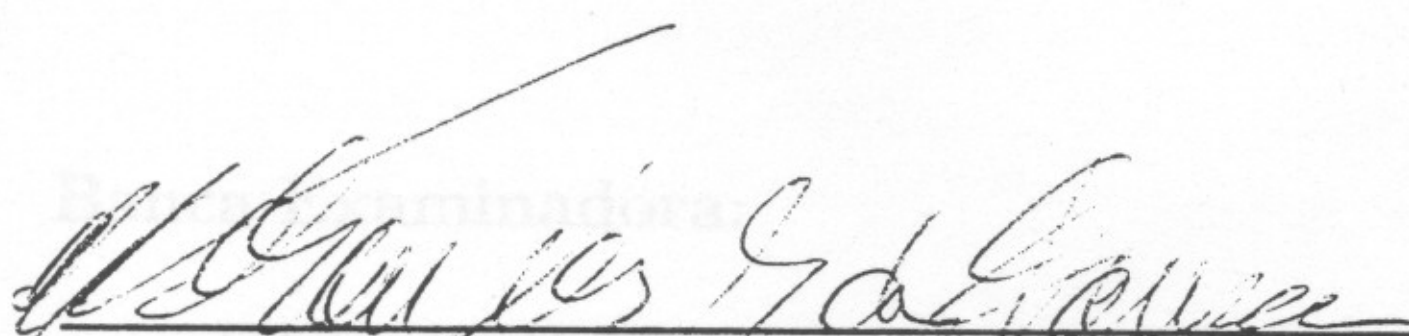
Programa de Pós-Graduação: Mestrado em Ciência da Computação

TERMO DE APROVAÇÃO

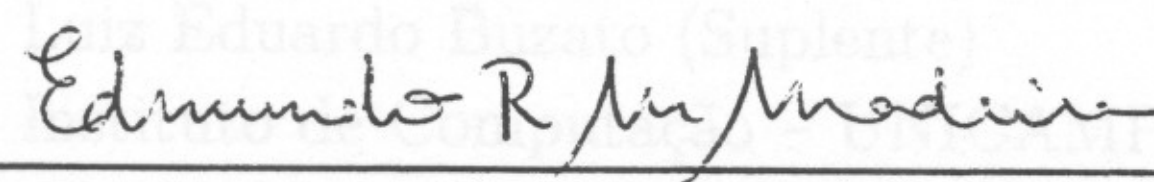
Dissertação defendida e aprovada em 29 de agosto de 2006, pela Banca examinadora composta pelos Professores Doutores:



Prof. Dr. Djamel Fawzi Hadj Sadok
Centro de Informática / UFPE.



Prof. Dr. Nelson Luís Saldanha da Fonseca
IC / UNICAMP.



Prof. Dr. Edmundo Roberto Mauro Madeira
IC / UNICAMP.

200723383

Gerência de Redes Virtuais Privadas de Camada 1

Neumar Costa Malheiros¹

Agosto de 2006

Banca Examinadora:

- Edmundo Roberto Mauro Madeira (Orientador)
- Djamel Fawzi Hadj Sadok
Centro de Informática – UFPE
- Nelson Luis Saldanha da Fonseca
Instituto de Computação – UNICAMP
- Luiz Eduardo Buzato (Suplente)
Instituto de Computação – UNICAMP

¹Suporte financeiro da CAPES.

Resumo

Um plano de controle distribuído, como a arquitetura GMPLS (*Generalized Multiprotocol Label Switching*), permite o provisionamento dinâmico de conexões em redes de transporte de camada 1, como redes ópticas ou redes TDM (*Time Division Multiplexing*). Dessa forma, essas redes podem oferecer serviços mais sofisticados como serviços VPN (*Virtual Private Network*). Esses serviços, então denominados VPN de camada 1 (L1VPN – *Layer 1 VPN*), permitem que a infra-estrutura de transporte do provedor seja compartilhada entre múltiplas redes clientes.

Neste trabalho, é proposta uma arquitetura para gerência de serviços L1VPN. A abordagem foi utilizar o paradigma de Gerência Baseada em Políticas (PBM – *Policy-Based Management*) para prover aos clientes um certo nível de controle e gerência sobre suas L1VPNs. Além disso, é apresentada a implementação de um protótipo da arquitetura proposta, assim como uma discussão das implicações de diferentes políticas para a gerência de configuração de serviços L1VPN, a partir de um estudo de caso.

Abstract

A distributed control plane architecture as GMPLS enhances transport networks with dynamic connection control. As a result, it allows the provisioning of advanced connectivity services, like Virtual Private Networks (VPNs), on layer 1 switching networks as optical and TDM networks. Such Layer 1 VPN (L1VPN) services enable multiple customer networks to share a single transport network.

In this work, an architecture for L1VPN management is proposed. The approach is based on Policy-Based Management (PBM) to provide customers with some level of control and management over their L1VPNs. Furthermore, a prototype implementation of the proposed architecture is presented and from a case study implications of different policies for L1VPN service configuration management are discussed.

Agradecimentos

Obrigado,

A todos os companheiros do IC,
principalmente, pelas idéias compartilhadas no “bandeijão”,
... e no “lava-jato”;

Ao pessoal do LRC, do LSD e ao Dr. Fábio Verdi,
pela paciência e amizade, pelo companheirismo no trabalho;

Ao Professor Edmundo, meu orientador,
pela oportunidade, pela confiança
e pelo profissionalismo e dedicação com que conduziu este trabalho;

Aos professores Nelson Fonseca e Djamel Sadok,
pelo tempo que dispensaram para que este trabalho fosse melhor;

Aos demais professores e profissionais do IC;

Aos meus padrinhos Marlon e Noélia,
A todos os meus familiares,
Aos meus amigos, em especial, Alessandra, Nayara, Atanásio e Rangel,
o apoio de todos vocês foi sempre muito importante;

Aos meus irmãos, Alex e Fernando,
pelo respeito e pela confiança;

Aos meus queridos pais, Neures e Mazinha,
meus primeiros *mestres*.

“Education’s purpose is to replace an empty mind with an open one.”

Sumário

Resumo	vii
Abstract	ix
Agradecimentos	xi
Lista de Figuras	xx
Lista de Acrônimos	xxi
I Introdução e Fundamentos	1
1 Introdução	3
1.1 Motivação	4
1.2 O Problema	5
1.3 Objetivos	5
1.4 Contribuições	6
1.5 Organização do texto	6
2 A Arquitetura GMPLS	9
2.1 Visão Geral	11
2.2 Gerência de Enlaces	13
2.3 Roteamento	14
2.4 Sinalização	15
3 Serviços L1VPN	19
3.1 Cenários de Aplicação	20
3.2 Modelo de referência	21
3.3 Alocação de recursos	23
3.4 Requisitos	24

3.5	Modelo Funcional	25
3.6	Tipos de Arquitetura	26
3.7	Modelos do Serviço	29
4	Gerência Baseada em Políticas	33
4.1	Visão Geral	36
4.2	Modelos de Informação de Políticas	37
4.3	<i>Framework</i>	39
4.4	Vantagens e Desvantagens	42
II	Proposta	45
5	Contexto do Trabalho	47
5.1	Gerência Baseada em Políticas	47
5.2	Serviços L1VPN	48
5.3	Contexto do Presente Trabalho	50
6	<i>Framework</i> de Políticas	53
6.1	Classes de Políticas	53
6.1.1	Políticas de Configuração	53
6.1.2	Políticas de Controle de Admissão	55
6.1.3	Políticas de Roteamento	56
6.2	<i>Framework</i>	57
6.2.1	Elemento de Gerência de Políticas	58
6.2.2	PDP e PEP	59
6.3	Representação de Políticas	59
7	Arquitetura Proposta	65
7.1	Modelo Funcional	65
7.2	Cenários de Aplicação	70
7.2.1	Abordagem Centralizada	71
7.2.2	Abordagem Distribuída	72
III	Avaliação e Conclusão	75
8	Implementação e Avaliação	77
8.1	Implementação	77
8.2	Estudo de Caso	81

8.2.1	Ambiente de Simulação	82
8.2.2	Primeiro Cenário	83
8.2.3	Segundo Cenário	85
8.2.4	Discussão Final	85
9	Conclusão	89
9.1	Trabalhos Futuros	91
	Bibliografia	93

Lista de Figuras

2.1	Arquitetura tradicional de uma rede de transporte.	9
2.2	Rede de transporte com plano de controle distribuído.	10
2.3	Encaminhamento MPLS.	12
2.4	Estabelecimento de conexão centralizado.	15
2.5	Estabelecimento de conexão distribuído.	16
2.6	Mensagens RSVP.	17
3.1	Modelo de referência do serviço L1VPN.	22
3.2	Esquema de identificação de um membro da VPN.	22
3.3	Modelo dedicado.	23
3.4	Modelo compartilhado.	23
3.5	Arquitetura centralizada.	26
3.6	Arquitetura distribuída.	27
3.7	Arquitetura híbrida.	28
3.8	Modelo de Serviço Baseado em Gerência.	30
3.9	Modelo de Serviço Baseado em Sinalização.	30
3.10	Modelo de Serviço de Sinalização e Roteamento.	31
4.1	Arquitetura de gerência de redes IP convencionais.	34
4.2	Níveis de gerência.	35
4.3	Visão geral da abordagem PBM.	36
4.4	Algumas classes do modelo PCIM.	39
4.5	<i>Framework</i> da abordagem PBM.	40
4.6	Cenário típico de uso do <i>framework</i> PBM.	41
4.7	Modelos de operação do protocolo COPS.	42
6.1	<i>Framework</i> baseado em políticas para gerência de Serviços L1VPN.	57
6.2	Modelo de informações de políticas para serviços L1VPN.	60
6.3	Exemplo de política em XML.	63
7.1	Arquitetura para gerência de serviços L1VPN.	66

7.2	Interações entre os módulos para estabelecimento de conexão.	69
7.3	Serviço L1VPN baseado em interface de gerência.	72
7.4	Serviço L1VPN baseado em interface de controle.	73
8.1	Estrutura do protótipo implementado.	77
8.2	Tela de gerência de membros.	79
8.3	Tela de controle de conexão.	80
8.4	Tela de edição de política.	80
8.5	Topologia de rede simulada.	81
8.6	Estrutura do ambiente de simulação.	82
8.7	Taxa de bloqueio de conexões.	84
8.8	Efeito de ativação da política.	85
8.9	Taxa de utilização de recursos da rede do provedor.	86
8.10	Taxa de bloqueio: modelo compartilhado x dedicado.	86

Lista de Acrônimos

ASTN *Automatically Switched Transport Network*

ATM *Asynchronous Transfer Mode*

BGP *Border Gateway Protocol*

CE *Customer Edge device*

COPS *Common Open Policy Service*

CPI *Customer Port Identifier*

CR-LDP *Constraint-based Label Distribution Protocol*

DTD *Document Type Definition*

DWDM *Dense Wavelength Division Multiplexing*

EMS *Element Management System*

FEC *Forwarding Equivalence Class*

FTP *File Transfer Protocol*

GMPLS *Generalized Multi-Protocol Label Switching*

HTML *HyperText Markup Language*

IETF *Internet Engineering Task Force*

IP *Internet Protocol*

IPsec *IP Security Protocol*

IS-IS *Intermediate System to Intermediate System*

ITU *International Telecommunication Union*

L1VPN *Layer 1 Virtual Private Network*

LMP *Link Management Protocol*

LSA *Link State Advertisement*

LSP *Label Switched Path*

LSR *Label Switching Router*

MIB *Management Information Base*

MPLS *Multiprotocol Label Switching*

NMS *Network Management System*

OAM *Operations, Administration and Maintenance*

OSPF *Open Shortest Path First*

OVPN *Optical VPN*

OXC *Optical Cross-Connect*

P *Provider device*

PBM *Policy-Based Management*

PCIM *Policy Core Information Model*

PDP *Policy Decision Point*

PE *Provider Edge device*

PEP *Policy Enforcement Point*

PMS *Provider Management System*

PPI *Provider Port Identifier*

QoS *Quality of Service*

RSVP *Resource Reservation Protocol*

RSVP-TE *Resource Reservation Protocol Traffic Engineering*

SC *Switched Connection*

SLA *Service Level Agreement*

SNMP *Simple Network Management Protocol*

SOAP *Simple Object Access Protocol*

SPC *Soft Permanent Connection*

SRLG *Shared Risk Link Group*

TCP *Transmission Control Protocol*

TDM *Time Division Multiplexing*

UNI *User-Network Interface*

VPN *Virtual Private Network*

XML *Extensible Markup Language*

XSLT *Extensible Stylesheet Language Transformations*

WDM *Wavelength Division Multiplexing*

Parte I

Introdução e Fundamentos

Capítulo 1

Introdução

A arquitetura das redes de transporte tradicionais foi direcionada pelo tráfego de voz. Em geral, essas redes oferecem apenas serviços estáticos de conexão ponto-a-ponto. O controle de tais conexões é realizado por sistemas de gerência centralizados e proprietários. Dessa forma, o provisionamento de serviços é lento e apresenta alto custo.

Essa arquitetura demonstrou-se ineficiente e sua operação e gerência bastante complexa devido a uma série de fatores como: o desenvolvimento de aplicações avançadas, com severos requisitos de qualidade de serviço; a convergência de redes; o aumento da demanda por capacidade de transmissão; e a necessidade de funções avançadas (dinâmicas) de admissão e provisionamento de conexões, agregação de tráfego e gerência de falhas.

Neste contexto, foi proposta a idéia de um plano de controle com uma arquitetura distribuída, primeiramente, com a especificação ASTN (*Automatically Switched Transport Network*), pela ITU (*International Telecommunication Union*). As principais funções de um plano de controle consistem no suporte ao provisionamento dinâmico¹ de conexões, descoberta automática de recursos e mecanismos eficientes para recuperação de conexões afetadas por falhas.

Uma proposta neste sentido é a arquitetura GMPLS (*Generalized Multi-Protocol Label Switching*) [31], definida pela IETF (*Internet Engineering Task Force*). A arquitetura GMPLS define mecanismos de sinalização e roteamento para o **controle dinâmico de conexões**, considerando várias tecnologias de comutação. Esses mecanismos são baseados principalmente em extensões de protocolos utilizados em redes IP.

Redes de transporte com um plano de controle como a arquitetura GMPLS podem estabelecer conexões de forma dinâmica para as redes clientes. Assim, redes de transporte de camada 1, como redes ópticas ou redes TDM (*Time Division Multiplexing*), podem oferecer serviços avançados de conectividade, como serviços de Redes Virtuais Privadas – VPN (*Virtual Private Network*).

¹Dinâmico significa de forma automatizada por meio de protocolos de rede.

Os serviços VPN, cujas conexões são estabelecidas em redes de camada 1, são então denominados **serviços VPN de Camada 1** (L1VPN – *Layer 1 VPN*) [45]. Serviços L1VPN permitem que vários clientes possam compartilhar a rede de transporte de um provedor. Esse serviço compreende um conjunto de funcionalidades que permitem um cliente interconectar suas redes, através do estabelecimento dinâmico de conexões de camada 1 na rede do provedor.

Um requisito fundamental é que deve ser oferecido ao cliente correspondente algum nível de controle e gerência sobre o serviço L1VPN. Além disso, existe o conceito de restrição de conectividade: conexões somente podem ser estabelecidas entre membros da mesma VPN. Outro aspecto é que informações de gerência e controle (como roteamento) devem ser separadas por VPN, de forma que um elemento de rede membro de uma VPN não deve receber informações de outra VPN.

1.1 Motivação

Neste contexto, estão relacionados a seguir alguns fatores que justificam a necessidade de um trabalho de pesquisa sobre gerência de serviços L1VPN:

Importância de serviços VPN: com os avanços nas tecnologias de redes de camada 1, como redes ópticas, e o desenvolvimento de arquiteturas de plano de controle baseadas na “inteligência” IP, serviços L1VPN representam uma solução flexível e eficiente para o suporte a múltiplos clientes. As vantagens do serviço L1VPN (como solução para compartilhamento de recursos de transporte e como modelo de negócios para as operadoras) e a tendência de arquiteturas de redes de transporte ópticas com plano de controle GMPLS motivam propostas de soluções para esses serviços. Existe uma forte expectativa de que L1VPN será um dos principais serviços das redes de próxima geração [44].

Sucesso de serviços VPN em redes MPLS: um dos grandes sucessos da arquitetura MPLS (*Multiprotocol Label Switching*) é o suporte a serviços VPN. A arquitetura GMPLS é uma evolução do MPLS. Com o avanço das tecnologias de redes ópticas com um plano de controle GMPLS, existe a necessidade de se definir mecanismos para o suporte a serviços VPN nessas novas arquiteturas de rede.

Ausência de propostas: a maioria dos trabalhos de padronização e de pesquisa sobre como prover serviços L1VPN estão voltados principalmente para aspectos do plano de controle. Não há propostas específicas para a gerência desses serviços.

1.2 O Problema

Em linhas gerais, o desafio em questão é como prover serviços L1VPN em redes de transporte com um plano de controle distribuído. Alguns dos problemas envolvidos são: endereçamento; isolamento de informações de roteamento por VPN; impacto sobre o mecanismo de controle de conexões; compartilhamento de informações sobre membros de VPNs, entre outros. Para lidar com os requisitos específicos desses serviços, podem ser necessários novos protocolos ou os existentes têm que ser modificados e aprimorados.

Em particular, este trabalho aborda o problema de **como gerenciar serviços L1VPN**, considerando principalmente os requisitos a seguir:

- O provedor deve oferecer a cada cliente certo nível de controle e gerência sobre seu serviço L1VPN;
- A gerência sobre a operação de cada VPN deve ser independente das demais;
- O cliente pode receber informações sobre a operação de sua VPN – *Operations, Administration and Maintenance (OAM) information*.

Assim, o problema tratado aqui é como múltiplos clientes, que compartilham a rede do provedor, podem gerenciar a operação de seus serviços L1VPN de forma independente. E ainda, como o provedor gerencia o nível de controle atribuído a cada cliente.

1.3 Objetivos

As questões de gerência de configuração de serviços L1VPN são o foco deste trabalho. O objetivo principal é propor uma **arquitetura para gerência de serviços L1VPN**. A abordagem considerada para lidar com os requisitos do problema descrito foi utilizar o paradigma de **Gerência Baseada em Políticas** (PBM – *Policy-Based Management*) [52]. Na abordagem PBM o administrador define um conjunto de regras de alto nível (as políticas) que controlam a utilização dos recursos e a operação dos serviços e da rede.

Os objetivos específicos deste trabalho são:

- Estudar a abordagem PBM, a arquitetura GMPLS e os conceitos e propostas para serviços L1VPN;
- Discutir como a abordagem PBM pode ser utilizada na gerência de serviços L1VPN;
- Propor uma arquitetura baseada em políticas para gerência de serviços L1VPN;
- Avaliar as implicações da abordagem PBM na gerência de serviços L1VPN.

O foco de interesse não é definir um conjunto específico de políticas para gerência de serviços L1VPN, mas sim discutir classes de políticas, como a arquitetura suporta o paradigma PBM para atender ao requisitos de gerência desses serviços e os efeitos da aplicação de diferentes políticas, por diferentes clientes que compartilham a rede do provedor.

1.4 Contribuições

A principal contribuição deste trabalho é a proposta de uma arquitetura baseada em políticas para a gerência de serviços L1VPN [30, 29]. No projeto da arquitetura proposta, considera-se que a rede do provedor possui um plano de controle distribuído, como a arquitetura GMPLS. A arquitetura define, sob a perspectiva do provedor, como cada cliente pode criar políticas para gerenciar seu serviço L1VPN. São apresentados cenários que caracterizam diferentes abordagens para aplicação da arquitetura proposta. Os cenários de aplicação descrevem como podem ser implementadas as funcionalidades do serviço L1VPN e os fatores que determinam o nível de controle que pode ser atribuído aos clientes do serviço.

Neste trabalho, é discutido como o *framework* de políticas da IETF pode ser utilizado no contexto de serviços L1VPN e também são propostas classes de políticas para gerência desses serviços. Além disso, são apresentados a implementação de um protótipo da arquitetura proposta e a avaliação das implicações do uso da abordagem PBM na gerência de serviços L1VPN, a partir de um estudo de caso.

1.5 Organização do texto

Esta dissertação está organizada em três partes: introdução e conceitos básicos; contexto e descrição da proposta; avaliação da proposta e conclusões. A seguir uma breve descrição do conteúdo deste documento:

Parte I: os demais capítulos desta primeira parte são uma revisão dos conceitos básicos sobre os temas envolvidos neste trabalho, conforme apresentado na literatura. O Capítulo 2 apresenta uma introdução à arquitetura GMPLS. O Capítulo 3 explica os conceitos, requisitos e *framework* de serviços L1VPN. Por fim, o Capítulo 4 introduz os fundamentos básicos sobre o paradigma de Gerência Baseada em Políticas.

Parte II: a proposta é descrita na segunda parte. Primeiramente, é definido o contexto do presente trabalho no Capítulo 5. O *framework* e as classes de políticas são discutidos no Capítulo 6. A descrição da arquitetura proposta aparece no Capítulo 7.

Parte III: na última parte, a implementação e avaliação da arquitetura são descritos no Capítulo 8 e as conclusões e trabalhos futuros são apresentados no Capítulo 9.

Capítulo 2

A Arquitetura GMPLS

Tradicionalmente, o desenvolvimento das redes de transporte era direcionado pelo tráfego de voz. A arquitetura dessas redes era organizada em *Plano de Dados* e *Plano de Gerência*. O Plano de Dados corresponde aos mecanismos e protocolos responsáveis pela transmissão dos dados entre os elementos de rede. O Plano de Gerência envolve os protocolos e sistemas para operação e gerência da rede.

Essas redes possuem uma **infra-estrutura de gerência centralizada**, composta de duas entidades principais: EMS (*Element Management System*) e NMS (*Network Management System*). O EMS é o sistema responsável pela gerência dos elementos de rede e é fornecido pelo fabricante do equipamento. O NMS é o sistema responsável pela gerência da rede. Tipicamente, ele é desenvolvido pelo próprio provedor. Em geral, ambos os sistemas utilizam tecnologias proprietárias e a operação da rede depende de intervenção manual (por exemplo, para o estabelecimento de conexões). A Figura 2.1 ilustra este cenário.

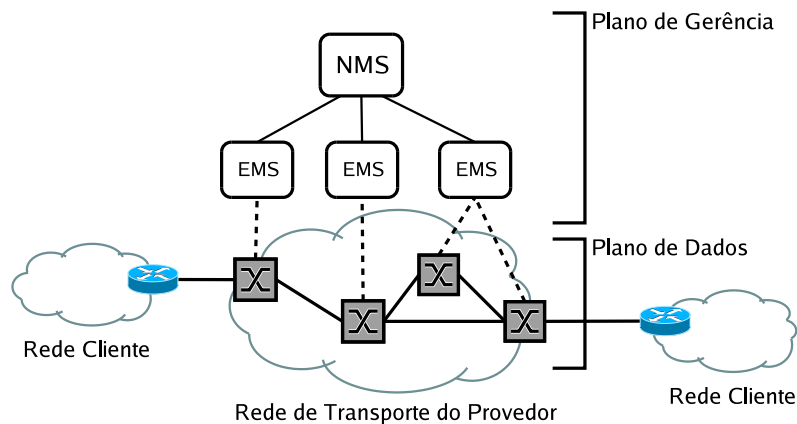


Figura 2.1: Arquitetura tradicional de uma rede de transporte.

O crescimento da Internet e o desenvolvimento de aplicações avançadas determinaram novas tendências. Atualmente, as redes de transporte dos provedores estão evoluindo para redes de comunicação de dados de alto desempenho. Os requisitos dos novos serviços e a dinâmica e o volume do tráfego de dados implicam em novos desafios para gerência e operação das redes de transporte atuais. A arquitetura das redes de transporte tradicionais, orientadas a tráfego estático de voz, demonstrou ser inadequada para este novo cenário. A sua infra-estrutura de gerência centralizada, baseada em sistemas proprietários semi-manuais, não suporta o provisionamento dinâmico de serviços, dificulta a interoperabilidade e a integração de novas tecnologias, torna a operação da rede mais lenta e propensa a erros, entre outros.

Diante desses problemas, ficou evidente a necessidade de uma nova arquitetura para as redes de transporte. Isso motivou a definição de um *Plano de Controle* com suporte a mecanismos dinâmicos e distribuídos que permitem um maior grau de automatização na gerência e operação da rede. O principal mecanismo do plano de controle é o controle automático de conexões. A natureza **descentralizada** das funções do plano de controle confere maior escalabilidade, robustez e eficiência aos mecanismos de operação da rede. Um plano de controle distribuído possibilita a automatização da operação da rede e assim, o suporte ao **aprovisionamento dinâmicos de serviços**. O primeiro passo importante no sentido de definir uma arquitetura de rede com plano de controle foi a especificação ASTN (*Automatically Switched Transport Network*), da ITU.

Essa nova abordagem é apresentada na Figura 2.2, onde o plano de controle é representado por um conjunto de entidades (lógicas) de controle associadas aos elementos de rede. Os mecanismos do plano de controle são implementados por meio de um conjunto de protocolos que suportam a comunicação de mensagens de controle entre essas entidades. A comunicação entre as entidades do plano de controle não precisa, necessariamente, utilizar os recursos da rede de dados. Ela pode ser realizada através de uma infra-estrutura de rede separada, dedicada para esse fim.

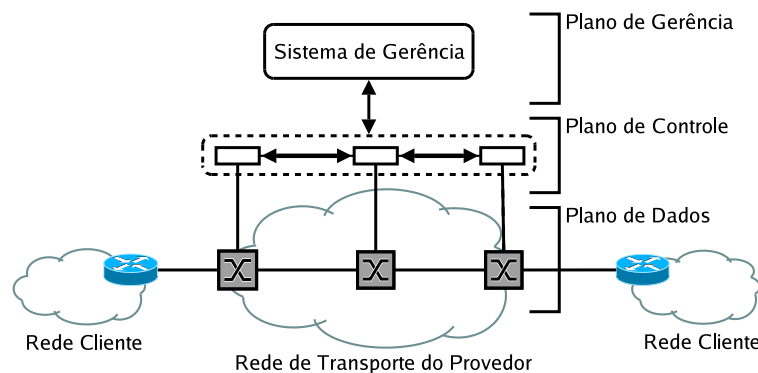


Figura 2.2: Rede de transporte com plano de controle distribuído.

Como apresentado em [6], as principais funções do plano de controle são:

Descoberta de elementos vizinhos (*neighbor discovery*): consiste em um elemento de rede identificar os outros elementos de rede aos quais ele está conectado no plano de dados, assim como as propriedades dos enlaces correspondentes;

Roteamento: esta função envolve dois procedimentos, a distribuição de informações de estado e da topologia da rede e o cálculo de rotas. Informações sobre os recursos e a topologia da rede são compartilhadas entre os elementos de rede através de um protocolo de roteamento. Estas informações são utilizadas no cálculo de rotas;

Sinalização: compreende os procedimentos para controle dinâmico de conexões. As entidades de controle comunicam entre si através de mensagens de sinalização para o estabelecimento automático de conexões;

Gerência de recursos locais: permite a uma entidade de controle gerenciar informações sobre o estado e disponibilidade dos recursos no elemento de rede correspondente.

2.1 Visão Geral

Existe uma expectativa de que a infra-estrutura das redes de próxima geração será constituída de redes óticas com um plano de controle distribuído, baseado em protocolos IP, interconectando roteadores de alto desempenho. Esta expectativa resulta do trabalho das principais organizações internacionais de padronização na especificação de um plano de controle para arquiteturas de redes de transporte. Neste contexto, pode-se destacar a arquitetura GMPLS (*Generalized Multi-Protocol Label Switching*), da IETF.

Como descrito em [31], a arquitetura GMPLS suporta o provisionamento dinâmico de conexões, assim como mecanismos de recuperação de conexões. O GMPLS é uma extensão da arquitetura MPLS [39]. A arquitetura MPLS define um **mecanismo de encaminhamento** para redes de comutação de pacotes¹ baseado em rótulos associados aos pacotes, o que permite a criação de rotas explícitas. De fato, essa arquitetura apresenta uma separação entre encaminhamento (plano de dados) e funções de plano de controle. A comunicação entre as entidades de controle pode utilizar uma infra-estrutura de rede separada. Como os protocolos do plano de controle MPLS requerem transporte IP, existe o conceito de um canal de controle IP.

A arquitetura MPLS apresenta várias vantagens em relação ao roteamento IP convencional, entre as quais pode-se destacar a simplificação do procedimento de encaminhamento de pacotes (o cabeçalho IP não precisa ser analisado em cada roteador) e o suporte a mecanismos de engenharia de tráfego.

¹Na verdade, rede de datagramas com comutação de pacotes ou células.

Em uma rede MPLS, um roteador é denominado LSR (*Label Switching Router*). A Figura 2.3 exemplifica como pacotes são encaminhados em uma rede MPLS. Quando um LSR de ingresso recebe um pacote IP, ele associa um rótulo ao pacote. No núcleo da rede, o encaminhamento do pacote é baseado no seu rótulo, que é comutado em cada roteador. O LSR de egresso remove o rótulo associado a um pacote antes de encaminhá-lo. No caso da figura, o comutador LSR1 associa o rótulo L1 ao pacote IP, o LSR2 troca o rótulo L1 por L2 e por fim, L2 é removido em LSR3.

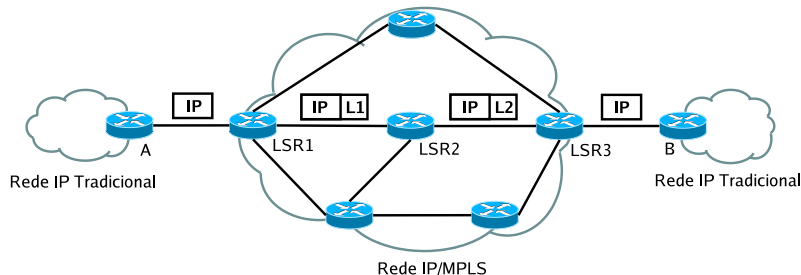


Figura 2.3: Encaminhamento MPLS.

Um rótulo é associado ao pacote de acordo com uma classificação baseada no seu cabeçalho IP, de forma que cada pacote pertence a uma classe denominada FEC (*Forwarding Equivalence Class*). Rotas explícitas podem ser definidas entre um par origem e destino. Neste caso, os rótulos associados ao pacote refletem o caminho dessa rota. Um protocolo de sinalização deve ser utilizado para distribuir os rótulos definidos para uma FEC por cada LSR, permitindo assim que os roteadores construam as tabelas de encaminhamento baseadas nos rótulos. A sequência de rótulos define um caminho em uma rede MPLS, que é então denominado LSP (*Label Switched Path*).

A arquitetura GMPLS estende o paradigma de controle MPLS para suporte a **várias tecnologias de comutação** através da generalização do conceito de rótulo. A arquitetura define uma **hierarquia de comutação** que inclui suporte a interfaces com diversas tecnologias de comutação: (1) comutação de pacotes; (2) comutação em camada 2 (circuitos virtuais); (3) comutação baseada em multiplexação no tempo; (4) comutação espacial baseada em multiplexação por comprimento de onda ou por fibra (em redes óticas). Assim, além de rótulos utilizados para associar pacotes a classes de tráfego, GMPLS suporta rótulos que identificam quadros de tempo, comprimentos de onda ou fibras, por onde os dados devem ser encaminhados.

De fato, os principais mecanismos da arquitetura MPLS são estendidos para o suporte a novas tecnologias de comutação. Pode-se destacar mecanismos de engenharia de tráfego [4] e de hierarquia de LSPs [22]. As principais características de engenharia de tráfego incluem a definição de rotas explícitas e o roteamento baseado em restrições. O

suporte a esses mecanismos exigiu o aprimoramento de protocolos de roteamento intra-domínio para suportar o compartilhamento de informações detalhadas do estado da rede.

O suporte a múltiplas tecnologias de comutação também exigiu novos mecanismos. Por exemplo, o uso da tecnologia DWDM (*Dense Wavelength Division Multiplexing*) levou a problemas de escalabilidade na propagação de informações de estado de enlace, devido ao grande número de enlaces (lógicos) entre dois elementos em uma rede ótica. Este problema motivou a técnica de agrupamento de enlaces (*link bundling*), na qual um conjunto de enlaces é tratado como se fosse um único enlace, do ponto de vista de engenharia de tráfego. Além disso, foi definido um novo protocolo para gerência de enlaces denominado *Link Management Protocol* (LMP).

A arquitetura GMPLS representa um **plano de controle comum** para redes com diferentes tecnologias de comutação no plano de dados. Em uma rede GMPLS, uma conexão (LSP) pode ser bidirecional e deve ser estabelecida entre interfaces compatíveis (com mesmo tipo de comutação). A separação entre plano de controle e plano de dados é um dos fatores que permitiu uma arquitetura de controle comum para redes de comutação de pacotes e redes de comutação de circuitos.

Em suma, mecanismos de sinalização e roteamento para o controle dinâmico de conexões na arquitetura GMPLS foram definidos a partir de: (1) extensões de protocolos utilizados em redes IP, como o OSPF (*Open Shortest Path First*) e o RSVP (*Resource Reservation Protocol*); (2) definição de um novo protocolo para gerência de enlace, o LMP. Portanto, como no MPLS, existe a necessidade de canais de controle IP. As subseções seguintes descrevem como as principais funções do plano de controle são implementadas no contexto da arquitetura GMPLS.

2.2 Gerência de Enlaces

Em redes IP convencionais, a descoberta de elementos vizinhos é realizada pelo protocolo de roteamento. Isso representa um problema no contexto da arquitetura GMPLS, pois em alguns casos, em redes óticas por exemplo, pode existir grande número de enlaces entre elementos de rede adjacentes. Por questões de **escalabilidade**, é inviável estabelecer adjacências de roteamento sobre vários enlaces entre o mesmo par de elementos de rede.

Assim, foi criado um protocolo específico para o mecanismo de gerência de enlaces, o LMP (*Link Management Protocol*) [26]. Este mecanismo permite a um elemento de rede descobrir seus nós vizinhos e informações sobre os enlaces entre ele e esses vizinhos. O LMP é um protocolo local, executado entre elementos de rede adjacentes no plano de dados. Ele permite identificar as propriedades dos enlaces entre esses elementos. As principais funções desempenhadas por este protocolo incluem [6]:

- Gerência da “conectividade” no plano de controle (estabelecimento, manutenção e gerência de canais de controle);
- Verificação de “conectividade” entre elementos;
- Correlação de propriedades de enlaces (verificação de compatibilidade dos parâmetros de configuração de enlaces);
- Gerência de falhas em enlaces (localização e notificação de falhas).

O protocolo LMP foi definido no contexto da arquitetura GMPLS, mas ele não depende dos outros mecanismos desta arquitetura, de maneira que ele pode ser utilizado em redes que não implementam o plano de controle GMPLS. Elementos de rede adjacentes trocam mensagens LMP através de um canal de controle entre eles.

2.3 Roteamento

O mecanismo de roteamento é responsável pela propagação de informações sobre recursos e topologia em um domínio de controle GMPLS. No entanto, esse mecanismo não define algoritmos de roteamento específicos a serem utilizados. Um conceito importante na arquitetura GMPLS é SRLG (*Shared Risk Link Group*), que define um grupo de enlaces que podem ser afetados por uma mesma falha. Informações dessa natureza são importantes para o cálculo de rotas disjuntas.

A arquitetura GMPLS também inclui o conceito de roteamento na origem (*source routing*). Neste caso, o cálculo de uma rota para uma conexão é realizado no nó de origem. Durante o estabelecimento da conexão, esta rota, denominada rota explícita, é repassada aos nós do caminho definido, utilizando-se um protocolo de sinalização.

Os protocolos de roteamento MPLS foram estendidos para atender aos requisitos da arquitetura GMPLS [24]. As novas tecnologias de comutação suportadas significam que novos atributos de enlaces precisam ser transportados pelos protocolos de roteamento. Além disso, na arquitetura GMPLS, enlaces agrupados (*bundled links*) podem ser tratados como um único enlace lógico do ponto de vista do roteamento. Isso reduz o volume de informações a serem compartilhadas.

Dois protocolos de roteamento intra-domínio foram adaptados para a arquitetura GMPLS, a saber, o OSPF (*Open Shortest Path First*) [23] e o IS-IS (*Intermediate System to Intermediate System*) [21]. No caso do OSPF, uma das extensões foi definir novos tipos de LSA (*Link State Advertisement*), para o suporte a novas propriedades de enlaces.

2.4 Sinalização

Como discutido em [6], o estabelecimento de uma conexão consiste basicamente em duas etapas: calcular uma rota para a conexão e configurar os elementos de rede correspondentes. A rota é calculada de acordo com o estado da rede e os parâmetros da requisição de conexão (como nós origem e destino, largura de banda, nível de proteção desejado, entre outros). Na segunda etapa, recursos são alocados e os comutadores são configurados de acordo com a rota definida.

A etapa de configuração, na qual um circuito é efetivamente estabelecido, pode ser realizada de forma centralizada ou distribuída. No primeiro caso, um sistema de gerência é responsável por configurar cada elemento de rede, como apresentado na Figura 2.4. O exemplo da figura representa o estabelecimento de uma conexão para interligar dois roteadores clientes, conectados aos nós A e C da rede do provedor de serviços. A rota é calculada pelo próprio sistema de gerência, que então configura cada elemento de rede no caminho para estabelecer a conexão. Neste caso, existe uma interface de gerência através da qual os elementos podem ser configurados.

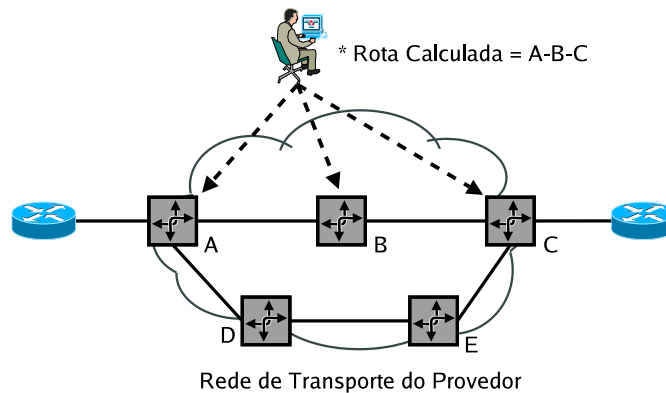


Figura 2.4: Estabelecimento de conexão centralizado.

Por outro lado, a **abordagem distribuída** é viabilizada por um mecanismo de sinalização do plano de controle. Mecanismos de sinalização representam um grande avanço diante da necessidade de automatizar o procedimento de estabelecimento de conexões em redes de comutação. Um mecanismo de sinalização envolve protocolos que suportam a troca de mensagens de controle entre os elementos de rede para o **estabelecimento dinâmico de conexões**. Esta tarefa pode envolver questões mais complexas que devem ser consideradas pelo mecanismo de sinalização. Entre elas pode-se destacar: conexões entre múltiplos domínios; modificação de atributos da conexão (como largura de banda); hierarquia de conexões; ou suporte à recuperação de conexões afetadas por falhas.

A Figura 2.5 ilustra o estabelecimento de uma conexão através do mecanismo de sinalização. Neste caso, uma requisição de conexão é enviada ao nó de origem. A conexão é efetivamente estabelecida através da troca de mensagens de controle entre os elementos de rede correspondentes (representada na figura pelas linhas tracejadas). Essas mensagens de sinalização controlam a alocação de recursos para estabelecimento da conexão.

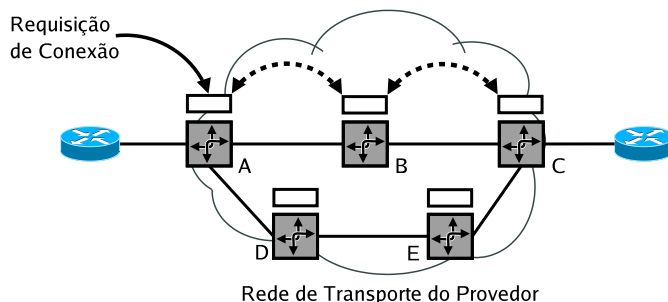


Figura 2.5: Estabelecimento de conexão distribuído.

A requisição de conexão pode ser enviada pelo sistema de gerência ou por um elemento da rede cliente. Uma conexão requisitada por gerência e configurada pelo plano de controle é denominada *Soft Permanent Connection* (SPC). Neste caso, a rota pode ser calculada pelo próprio sistema de gerência e repassada na requisição ao nó de origem. Ou então, a rota pode ser calculada pelo nó que recebe a requisição (*source routing*). Por outro lado, se a requisição é enviada por um elemento da rede cliente, através de uma interface do plano de controle, a conexão é denominada *Switched Connection* (SC).

O mecanismo de sinalização da arquitetura GMPLS [5] permite aos elementos de rede requisitar ou alocar rótulos para uma conexão, informar aos vizinhos sobre os rótulos alocados e realizar as configurações de comutação apropriadas para encaminhar os dados. A alocação de rótulos em um comutador é uma decisão local (independe dos rótulos alocados em outros comutadores).

A sinalização GMPLS considera dois protocolos: RSVP-TE (*Resource Reservation Protocol Traffic Engineering*) [4] e CR-LDP (*Constraint-based Label Distribution Protocol*). Estes protocolos surgiram no contexto da arquitetura MPLS e foram então estendidos para a arquitetura GMPLS. As funcionalidades incorporadas a esses protocolos foram delineadas principalmente em função do suporte a múltiplas tecnologias de comutação. As principais características adicionadas são: (1) generalização do conceito de rótulo (um rótulo pode identificar um *slot* de tempo, um comprimento de onda, etc); (2) suporte a um rótulo comum para todos os segmentos de uma conexão; (3) suporte a conexões bidirecionais; (4) separação (física) entre plano de controle e de dados (conexões estabelecidas não são afetadas por falhas no plano de controle); (5) procedimentos para recuperação de estado (do controle de sinalização).

Entre os dois protocolos citados, o RSVP-TE tem mais destaque na literatura. Este protocolo inclui mensagens para requisição (*Path*) e confirmação da alocação (*Resv*) de rótulos (e recursos) e também suporte à especificação de rotas explícitas e de parâmetros de requisição de conexões. Ele permite ainda especificar associações entre conexões (LSPs) relacionadas.

A Figura 2.6 exemplifica o estabelecimento de uma conexão entre os nós A e C. As mensagens de sinalização não são enviadas junto com os dados (como na arquitetura MPLS). Elas são enviadas através de um canal de controle IP no plano de controle GMPLS, independente da rede de dados. As mensagens *Path* contêm parâmetros da requisição de conexão (como largura de banda). As mensagens *Resv* confirmam a alocação de recursos e o rótulo então associado à conexão. Quando do estabelecimento de conexões bidirecionais, as mensagens *Path* já informam para o nó seguinte no caminho, o rótulo alocado para o sentido de volta da conexão.

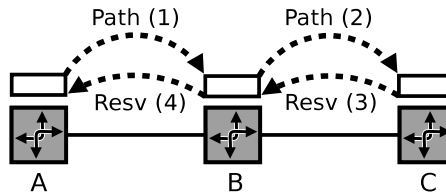


Figura 2.6: Mensagens RSVP.

O mecanismo de sinalização também é utilizado em procedimentos de **recuperação de conexões** afetadas por falhas. Em geral, estes procedimentos consistem em rotear o tráfego afetado para rotas alternativas. Na arquitetura GMPLS existem dois tipos de recuperação: **proteção** (*protection*) e **restauração** (*restoration*) [32]. No primeiro caso, recursos são previamente alocados e então dedicados para serem utilizados na recuperação de conexões específicas, ou seja, conexões secundárias previamente estabelecidas ficam disponíveis para substituir conexões ativas afetadas. Por outro lado, no procedimento de restauração, as conexões alternativas são estabelecidas de forma dinâmica mediante ocorrência de falha. A vantagem do primeiro método é o menor tempo de recuperação, enquanto o segundo possibilita melhor utilização dos recursos da rede.

O mecanismo de sinalização desempenha três funções básicas envolvidas na recuperação de conexões:

- Estabelecimento da conexão alternativa utilizada para substituir a conexão afetada;
- Notificação da ocorrência de falhas por meio de mensagens de controle específicas;
- Requisição de desvio do tráfego de uma conexão afetada para a conexão alternativa.

Capítulo 3

Serviços L1VPN

Um serviço VPN utiliza uma infra-estrutura de rede compartilhada para prover comunicação segura e restrita a um grupo de entidades geograficamente dispersas. Tradicionalmente, este serviço utiliza o conceito de “conexões virtuais” de camada 3 ou 2 para oferecer conectividade entre duas redes ou entre um hospedeiro e uma rede.

O controle dinâmico de conexões em redes de comutação de circuitos, através de um plano de controle distribuído, motivou a idéia de utilizar conexões de camada 1 para prover serviços VPN. Esses serviços são então denominados **VPN de Camada 1** (L1VPN). Dessa forma, serviços VPN podem ser oferecidos através de conexões na camada física, por exemplo, em redes TDM ou em redes ópticas com multiplexação por comprimento de onda (WDM – *Wavelength Division Multiplexing*). Neste caso, o serviço L1VPN também é chamado VPN óptica (OVPN – *Optical VPN*).

O serviço L1VPN compreende um conjunto de funcionalidades que suporta a **alocação dinâmica de recursos** da rede do provedor. Assim, um cliente pode interconectar suas diversas redes através do estabelecimento de conexões de camada 1 na rede do provedor, de acordo com uma topologia desejada e com sua demanda de tráfego. Dessa forma, serviços L1VPN possibilitam que vários clientes possam **compartilhar** a rede de transporte de um provedor, de maneira que não é necessário que cada cliente assuma os custos de implantar sua própria rede de transporte. O trabalho de operação e manutenção da infra-estrutura de transporte fica a cargo do provedor e os custos podem ser divididos entre os vários clientes.

Serviços L1VPN são diferentes de outros serviços de camada 1 em dois aspectos principais. As conexões VPN podem ser estabelecidas somente entre um conjunto bem definido de entidades, que são os membros da VPN (*VPN membership*). Segundo, o provedor deve oferecer, de forma independente, a cada cliente de serviços L1VPN, algum nível de gerência e controle, incluindo a capacidade de alterar a topologia da VPN. Além disso, o serviço VPN é caracterizado pelo isolamento das informações de controle de cada VPN.

As principais vantagens desse serviço são discutidas em [44]. Os clientes podem expandir sua rede a partir da alocação de recursos de rede de transporte de alto desempenho, sem a necessidade de investir em uma infra-estrutura de rede própria. Os recursos podem ser alocados dinamicamente, de acordo com suas necessidades, para atender às alterações na demanda de tráfego. A operação da rede de transporte é delegada ao provedor e os custos são compartilhados. Os provedores podem oferecer serviços de conectividade avançados, com novas oportunidades de negócios.

Muitos trabalhos de pesquisa e esforços de padronização estão sendo realizados a fim de que se possa explorar todo o potencial de serviços L1VPN. A maioria das iniciativas é direcionada a questões relacionadas com o plano de controle, com novas propostas ou extensões de protocolos e mecanismos para o suporte a esses serviços. Muitas questões ainda precisam ser resolvidas neste contexto, assim como na área de gerência. Por exemplo, como prover que cada cliente possa gerenciar sua VPN de forma independente. Pois, em geral, a rede de transporte é compartilhada entre diferentes tipos de redes de serviços, como redes IP e ATM (*Asynchronous Transfer Mode*). Cada uma dessas redes clientes pode ser administrada por diferentes organizações ou diferentes departamentos de uma organização, o que resulta em diferentes políticas operacionais e estratégias de gerência.

3.1 Cenários de Aplicação

Serviços L1VPN são apropriados para clientes com grandes redes corporativas ou provedores de pequeno e médio porte com alta demanda por largura de banda, que desejam interligar suas redes através de uma infra-estrutura de transporte de alto desempenho. Esse serviço representa uma solução flexível para um provedor compartilhar sua rede de transporte entre vários clientes, permitindo a alocação de recursos de camada 1 sob demanda. O trabalho apresentado em [42], em desenvolvimento na IETF, discute cenários de aplicação para serviços L1VPN, como descrito a seguir:

Multi-Service Backbone: neste cenário, um provedor utiliza serviços L1VPN para compartilhar sua rede entre seus diversos departamentos que, em camadas superiores, oferecem diferentes redes de serviço, como redes IP, ATM ou MPLS. Cada departamento aloca os recursos de acordo com suas próprias necessidades e as características do serviço que oferece. Assim, a infra-estrutura de transmissão do provedor é compartilhada entre vários serviços, de forma flexível e com funções de controle independentes para cada departamento.

Carrier's Carrier: neste caso, um provedor oferece serviço L1VPN para outro provedor, que por sua vez, oferece serviços para redes clientes. Neste caso, se os provedores pertencem a diferentes organizações, as informações e funções de controle do serviço

podem ser mais limitadas do que no cenário anterior. Por exemplo, o serviço L1VPN pode oferecer apenas uma visão limitada da topologia da rede de transporte. O provedor que recebe serviço L1VPN pode se concentrar nos serviços específicos que ele oferece para as redes clientes, pois a infra-estrutura de comunicação de camada 1 é responsabilidade do provedor de serviço L1VPN.

Negociação de recursos: este cenário é caracterizado por um cliente que pode receber serviços L1VPN de mais de um provedor. Neste caso, o cliente pode estabelecer conexões em diferentes provedores para alcançar um mesmo destino. As vantagens incluem redundância de rotas que pode ser usada como mecanismo para recuperação de falhas e aumento do nível de disponibilidade do serviço. Além disso, o cliente tem mais flexibilidade para escolher entre diferentes provedores, considerando os custos e as facilidades dos serviços oferecidos.

Escalonamento de recursos: em algumas aplicações específicas, como *backup* de dados, um cliente L1VPN precisa alocar recursos do provedor em períodos bem definidos e previamente determinados. Neste cenário, o provedor deve suportar mecanismos de escalonamento que permitam programar a reserva de recursos e o estabelecimento de conexões. Programar previamente um período para alocação e liberação dos recursos pode contribuir para uma melhor utilização dos recursos.

3.2 Modelo de referência

A Figura 3.1 apresenta o modelo de referência do serviço L1VPN [45]. Este modelo define a **interface de serviço L1VPN**, através da qual os clientes acessam as funcionalidades do serviço. Essa interface pode ser implementada no plano de controle ou através de um sistema de gerência. Além disso, o modelo define caracteriza três tipos de elementos de rede:

- CE (*Customer Edge device*): é um elemento de borda da rede do cliente que está conectado a pelo menos um elemento do provedor. O CE provê acesso ao serviço L1VPN. Ele pode ser, por exemplo, um roteador ou um comutador de camada 2.
- PE (*Provider Edge device*): é um elemento de borda da rede do provedor ao qual pelo menos um CE está conectado. O PE provê serviços L1VPN para o CE. Um PE pode ser qualquer elemento com capacidade de comutação em camada 1, como um comutador óptico OXC (*Optical Cross-Connect*) ou TDM.
- P (*Provider device*): é um elemento do núcleo da rede do provedor que não está conectado a nós de redes clientes. É um comutador de camada 1, como o PE.

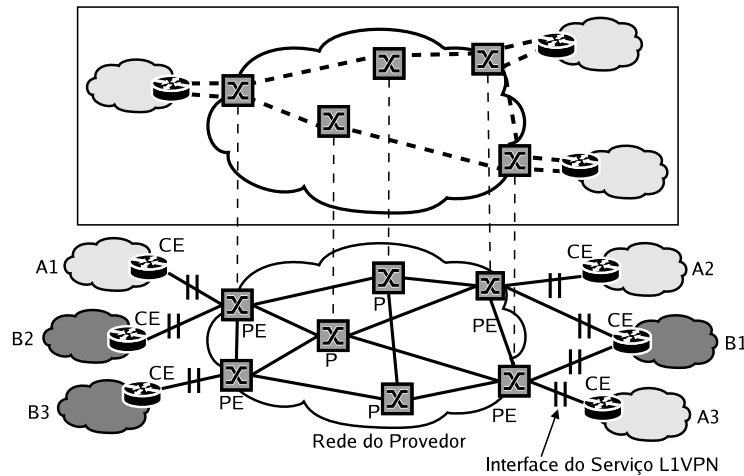


Figura 3.1: Modelo de referência do serviço L1VPN.

Uma conexão VPN é estabelecida entre dois CEs. Conexões são permitidas somente entre CEs membros da mesma VPN. Portanto, um CE que pertence a um cliente não pode estabelecer uma conexão com um CE que pertence a outro cliente. No cenário apresentado na figura, a rede de transporte do provedor é compartilhada por dois clientes (A e B). A parte superior da figura mostra uma possível topologia para interconectar as redes do cliente A. As conexões da L1VPN são representadas pelas linhas tracejadas.

Um membro de uma VPN é designado por um par de portas (lógicas) que representam os extremos de uma conexão (lógica e estática) entre um CE e um PE no plano de dados. Esse par é formado por dois identificadores: o identificador de porta do cliente CPI (*Customer Port Identifier*) e o identificador de porta do provedor PPI (*Provider Port Identifier*). Portanto, um par CPI-PPI identifica um membro da VPN. Além disso, pode existir uma conexão entre CE e PE no plano de controle, através da qual são acessadas as funcionalidades do serviço L1VPN. Deve haver um esquema para separar as mensagens de controle das VPNs. Esses aspectos estão ilustrados na Figura 3.2.

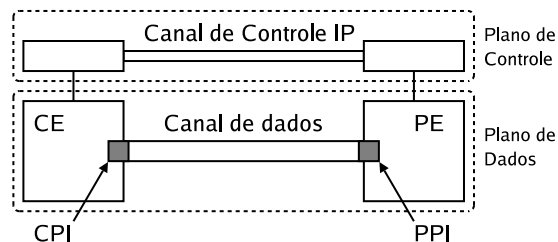


Figura 3.2: Esquema de identificação de um membro da VPN.

3.3 Alocação de recursos

Um aspecto muito importante no suporte a serviços L1VPN é como o provedor gerencia a alocação de recursos. Como explicado em [45], existem dois modelos de alocação de recursos: **dedicado** (*dedicated*) e **compartilhado** (*shared*).

No modelo de alocação dedicado, os recursos da rede do provedor são **reservados exclusivamente** para uma L1VPN específica. Os recursos dedicados para uma L1VPN não podem ser alocados para nenhuma outra L1VPN, mesmo que não estejam sendo utilizados. Neste modelo, a rede do provedor é particionada através da reserva dos recursos para os serviços L1VPN. A Figura 3.3 ilustra este cenário. Como os recursos são reservados, o provedor pode fornecer informações detalhadas de uma porção da rede para o cliente correspondente, de maneira que este cliente tem condições de realizar funções mais avançadas de roteamento e otimização da utilização de recursos.

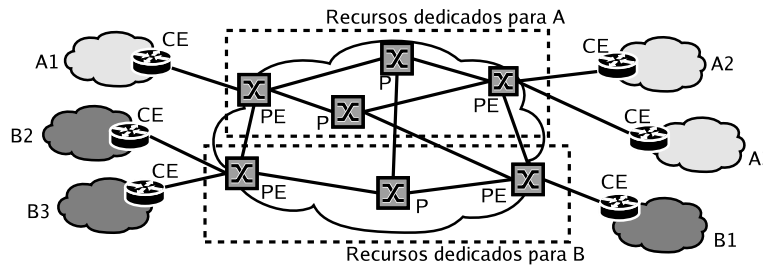


Figura 3.3: Modelo dedicado.

No modelo de alocação compartilhado, os recursos da rede do provedor são **compartilhados no tempo** por múltiplos clientes. Assim, um recurso pode ser alocado para qualquer uma, dentre as L1VPNs que compartilham aquele recurso. Após liberado, aquele recurso pode ser alocado para outra L1VPN. Do ponto de vista do cliente, este modelo limita algumas funcionalidades, como otimização de alocação de banda ou cálculo de rotas explícitas, pois como os recursos são compartilhados o cliente não tem informações detalhadas da rede do provedor. A Figura 3.4 ilustra este cenário.



Figura 3.4: Modelo compartilhado.

3.4 Requisitos

Além das características discutidas anteriormente, uma série de requisitos foram definidos para serviços L1VPN [45]. Os principais desses requisitos são brevemente apresentados a seguir:

- O provedor deve oferecer aos clientes controle dinâmico de conexões. Assim, os clientes podem estabelecer, remover ou modificar conexões na rede de transporte. Sob esse ponto de vista, a topologia da VPN está sobre controle do cliente. No entanto, o controle de conexão está sujeito a algumas restrições (impostas pelo provedor), como limitação de recursos, permissões, dentre outras.
- Clientes devem ser notificados quando uma requisição de conexão é rejeitada.
- O cliente deve ter acesso às informações sobre os membros de sua VPN. Isso permite que um CE possa identificar a quais outros CEs ele pode se conectar, ou seja, quais outros CEs pertencem à mesma VPN.
- O provedor deve suportar mecanismos de controle de autorização e autenticação de clientes.
- O provedor pode fornecer ao cliente informações sobre disponibilidade de recursos e topologia da rede de transporte. Desta maneira, o cliente pode determinar se uma conexão pode ser estabelecida entre dois elementos. Mais importante, ele tem condições de projetar a topologia de sua VPN e implementar funções de engenharia de tráfego. Por questões como escalabilidade, essas informações podem ser resumidas ou agregadas.
- O cliente pode também receber informações sobre disponibilidade de recursos e topologia de suas redes remotas.
- Conexões são restritas aos membros de uma mesma VPN (*connectivity restriction*). O provedor deve verificar se destino e origem de uma conexão requisitada são membros da mesma VPN, caso contrário, a conexão não deve ser estabelecida.
- O provedor deve ser capaz de manter, de forma independente, informações de membros e topologia de cada VPN, para que essas informações possam ser acessadas pelos clientes correspondentes.
- Clientes podem utilizar endereços privados em suas VPNs.
- O provedor deve oferecer ao cliente certo nível de controle e gerência sobre sua VPN, incluindo a capacidade de reconfigurar sua topologia.

- Um cliente pode especificar políticas de operação de sua VPN.
- O provedor deve refletir as políticas operacionais dos diferentes clientes nas respectivas VPNs de forma independente.
- O provedor pode fornecer informações sobre operação e administração (*OAM information*) de uma VPN para o cliente correspondente, isso inclui informações sobre desempenho e ocorrência de falhas.

Alguns requisitos não são obrigatórios, enquanto outros dependem dos mecanismos e modelos utilizados para prover o serviço. Por exemplo, informações sobre disponibilidade de recursos da rede do provedor são repassadas ao cliente somente se o modelo de alocação de recursos é dedicado. Os quatro últimos requisitos estão relacionados com a gerência de serviços L1VPN e, portanto, são os mais importantes no contexto deste trabalho.

3.5 Modelo Funcional

O modelo funcional do serviço L1VPN é apresentado em [45]. São discutidas quatro categorias principais de funções necessárias para o provisionamento desses serviços:

Controle de informações de membros: compreende as funções para gerenciar as informações sobre os membros de uma VPN. Conexões somente podem ser estabelecidas entre membros da mesma VPN, mas podem haver restrições adicionais mesmo entre os próprios membros. As informações sobre membros podem ser configuradas de forma estática ou a rede do provedor pode suportar um mecanismo de descoberta automática (*VPN membership autodiscovery*).

Funções de roteamento: consiste principalmente no compartilhamento de informações de roteamento e cálculo de rotas. O provedor pode gerenciar três tipos de informações de roteamento: topologia e estado de sua própria rede, da rede do cliente e informações sobre as conexões (*connectivity information*). Essas informações, juntamente com outras restrições e configurações, são utilizadas no cálculo de rotas para as conexões requisitadas. O cálculo de rota também depende do modelo de alocação de recursos.

Controle de conexões: envolve o conjunto de mecanismos utilizados para estabelecer ou remover conexões, como os mecanismos de sinalização do plano de controle. Uma conexão é estabelecida de acordo com uma rota explícita, previamente calculada, ou sua rota pode ser definida de forma dinâmica, a partir de decisões locais, em cada nó (*hop-by-hop*). Pode haver a necessidade de mapeamento de endereços, no caso da VPN utilizar endereços privados.

Funções de gerência: consiste em funções tradicionais de gerência, como gerência de configuração, desempenho e falhas, mas com uma particularidade: as ações de gerência para cada serviço L1VPN são independentes, assim eventos e decisões de gerência devem ser tratados de acordo com a VPN à qual eles estão relacionados. Outros aspectos de gerência incluem mecanismos de autenticação, autorização e contabilidade (*accounting*). Além disso, o provedor deve oferecer suporte para que clientes possam especificar políticas de gerência. O provedor também deve gerenciar e disponibilizar para os clientes, informações administrativas e operacionais (OAM *information*) para cada VPN, assim como suportar funções para gerência de falhas.

É importante observar que apenas as funções relacionadas com a primeira categoria são específicas para serviços L1VPN. No entanto, as funções de roteamento, controle de conexões e gerência, já presentes na rede do provedor, precisam ser estendidas para suportar esses serviços.

3.6 Tipos de Arquitetura

Conforme descrito em [45, 19], a implementação do modelo funcional do serviço L1VPN pode seguir três tipos de arquitetura: centralizada, distribuída ou híbrida. O tipo de arquitetura é caracterizado de acordo com o local onde são implementadas as funcionalidades do serviço. Elas podem ser implementadas em elementos de rede específicos (CE, PE ou P) ou no sistema de gerência da rede do provedor.

Na **arquitetura centralizada** existe a figura de um sistema de gerência de rede do provedor PMS (*Provider Management System*). As funcionalidades do serviço L1VPN são implementadas no PMS. Mecanismos distribuídos de roteamento ou sinalização não são utilizados. A Figura 3.5 demonstra a arquitetura centralizada.

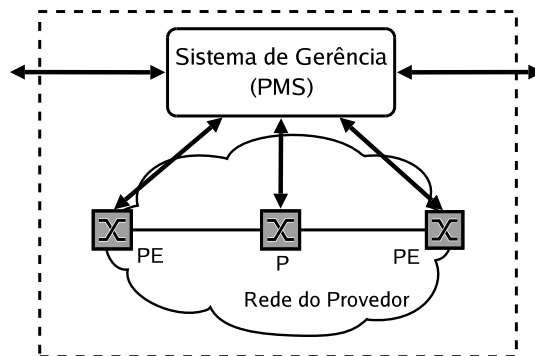


Figura 3.5: Arquitetura centralizada.

Nesta arquitetura, existe uma interface de gerência entre cliente e provedor, através da qual o cliente acessa as funcionalidades do serviço. Por exemplo, clientes podem acessar o PMS para obter informações sobre membros de VPNs ou sobre topologia e disponibilidade de recursos ou para requisitar conexões. O PMS comunica com os elementos de rede para obter informações de roteamento. Cada elemento possui apenas informações de roteamento locais, eles não precisam reter informações sobre a topologia completa da rede. De forma similar, o PMS também comunica com os elementos de rede para o estabelecimento de conexões, sendo necessário configurar cada elemento na rota definida. O cálculo de rotas também é realizado no PMS.

Na **arquitetura distribuída** as funcionalidades do serviço L1VPN são implementadas no elementos P e PE. Também são utilizados mecanismos distribuídos de roteamento e controle de conexões. A comunicação entre o cliente e o provedor é realizada por meio de mensagens no plano de controle. Esta arquitetura é ilustrada na Figura 3.6.

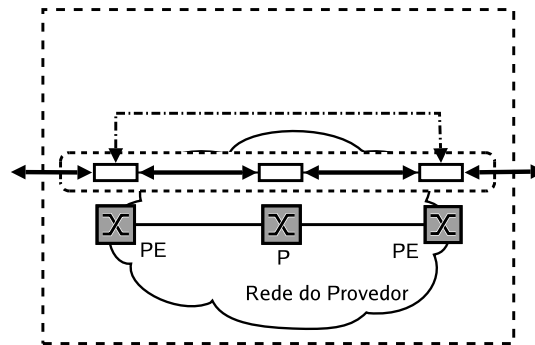


Figura 3.6: Arquitetura distribuída.

Neste caso, as informações sobre membros de VPNs são mantidas nos elementos de borda (PEs). Estes elementos comunicam entre si para compartilhar essas informações, assim como podem enviar para um CE as informações relacionadas com a VPN da qual ele faz parte. Assim, um PE apenas mantém informações sobre os membros de uma VPN se existe algum membro desta VPN conectado a ele.

Os elementos de borda também são os únicos que mantêm informações sobre a rede do cliente. Eles comunicam entre si e com os CEs locais para compartilhar essas informações. Por outro lado, tanto PEs quanto Ps devem manter informações de roteamento da rede do provedor. No caso do modelo de alocação dedicado, um PE pode fornecer essas informações (mesmo que abstratas) aos CEs a ele conectados. Dessa forma, o cálculo de rotas pode ser realizado pelo CE ou PE (*source routing*).

Os elementos PE e P comunicam entre si, através de um mecanismo de sinalização, para estabelecer ou remover conexões. Além disso, o PE é responsável por processar requisições de conexão enviadas por um CE através de uma interface no plano de controle.

A **arquitetura híbrida** combina mecanismos centralizados e distribuídos. Em geral, funções específicas de serviços L1VPN, como gerência de informações de membros, são centralizadas. Enquanto que funções como roteamento e controle de conexões podem ser distribuídas. A Figura 3.7 demonstra esta arquitetura. O cliente pode acessar algumas funcionalidades do serviço através do plano de controle e outras através da entidade de gerência (PMS).

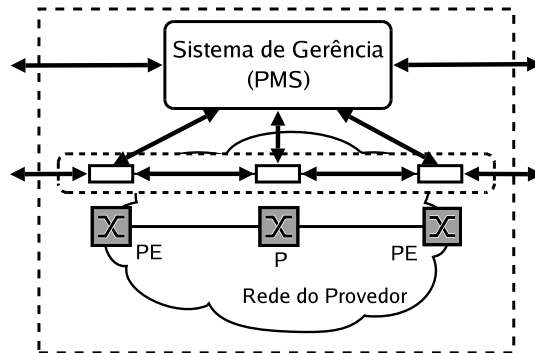


Figura 3.7: Arquitetura híbrida.

O trabalho apresentado em [47] apresenta uma comparação entre as vantagens e desvantagens das diferentes arquiteturas. A arquitetura centralizada representa um solução de curto prazo pois ela facilita a implementação e a integração com sistemas legados, mas é deficiente em aspectos de **escalabilidade** e **desempenho**. Esses fatores são discutidos a seguir:

- Facilidade de implementação: funções podem ser implementadas no próprio Sistema de Gerência de Rede do provedor NMS (*Network Management System*); protocolos de gerência existentes (como o SNMP) podem ser utilizados para configuração dos elementos de rede; não há necessidade de implementação de mecanismos distribuídos do plano de controle.
- A integração com sistemas legados também resulta do fato de que os elementos de rede não precisam implementar os protocolos do plano de controle.
- A necessidade de troca de informações de uma entidade centralizada com o cliente e com os elementos de rede através de uma interface de gerência é lenta compromete o desempenho de operações de estabelecimento de conexões e notificação e recuperação de falhas.
- Uma entidade centralizada representa um ponto único de falha o que afeta a robustez do sistema.

- Para coletar informações de roteamento, a entidade de gerência precisa se comunicar com cada elemento da rede e com cada elemento na rota calculada, para estabelecer uma conexão. Essa entidade ainda precisa manter informações de estado para controle de conexão e roteamento. Todos esses fatores comprometem a escalabilidade do sistema.

Por outro lado, a arquitetura distribuída emprega mecanismos distribuídos de sinalização e roteamento, o que aumenta a escalabilidade e o desempenho do sistema. No entanto, ela exige a implementação de um plano de controle distribuído, o que ainda representa muitos desafios.

No que diz respeito à arquitetura híbrida, ele representa uma alternativa flexível que permite combinar as vantagens das outras arquiteturas. Neste caso, um provedor pode implementar a arquitetura centralizada e, progressivamente, agregar mecanismos distribuídos. Funcionalidades específicas do serviço L1VPN podem ser mantidas de forma centralizada, enquanto que funções comuns, como mecanismos para estabelecimento de conexões, podem então ser distribuídos.

3.7 Modelos do Serviço

A IETF define três modelos de serviço L1VPN [42, 44]: modelo baseado em gerência, modelo baseado em sinalização e modelo baseado em sinalização e roteamento. A definição dos modelos é baseada na interface do serviço e na semântica das mensagens na comunicação entre cliente e provedor. Cada modelo pode oferecer diferentes funcionalidades e métodos de acesso. A implementação de um determinado modelo depende de vários fatores como o cenário de aplicação, os requisitos do cliente, o contrato de nível de serviço SLA (*Service Level Agreement*) e as tecnologias de rede do provedor e do cliente, principalmente no que diz respeito aos mecanismos do plano de controle.

No **modelo baseado em gerência**, as funções do serviço L1VPN são implementadas como parte do sistema de gerência do provedor. Neste caso, as redes clientes acessam essas funções através de uma **interface de gerência**. Não há transferência de mensagens de plano de controle entre cliente e provedor. Portanto, o cliente (CE) não precisa implementar os protocolos correspondentes. A Figura 3.8 demonstra este modelo.

O cliente pode acessar o sistema de gerência para obter informações de falha e desempenho de sua VPN ou para enviar uma requisição de conexão. Neste caso, o sistema então se encarrega de estabelecer uma conexão de camada 1 através da sua rede de transporte, para interconectar os CEs correspondentes. Para tal, ele comunica com o plano de controle, enviando uma requisição de estabelecimento da conexão ao respectivo nó de ingresso (o PE ao qual o CE de origem está conectado, conforme a requisição do cliente).

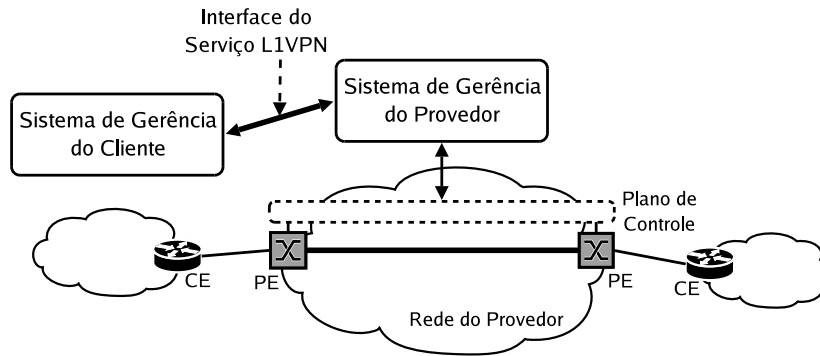


Figura 3.8: Modelo de Serviço Baseado em Gerência.

O estabelecimento da conexão no núcleo da rede do provedor é efetuado pelo mecanismo de sinalização do plano de controle. Tal conexão, acionada por um sistema de gerência e estabelecida pelo plano de controle, é denominada *Soft Permanent Connection (SPC)*. Dessa forma, os CEs podem estabelecer uma adjacência de roteamento utilizando a conexão estabelecida entre os respectivos PEs. Assim, esse modelo caracteriza um cenário *overlay*. Ele pode ser implementado segundo uma arquitetura centralizada ou híbrida.

No **modelo baseado em sinalização**, o cliente acessa as funções do serviço L1VPN através de uma **interface no plano de controle**. Entretanto, as mensagens de controle se restringem ao contexto de mecanismos de sinalização. Este modelo também caracteriza um cenário *overlay*. O cliente ainda pode obter informações de gerência a partir do sistema de gerência do provedor. Este modelo pode ser implementado segundo uma arquitetura distribuída ou híbrida. A Figura 3.9 ilustra este modelo de serviço.

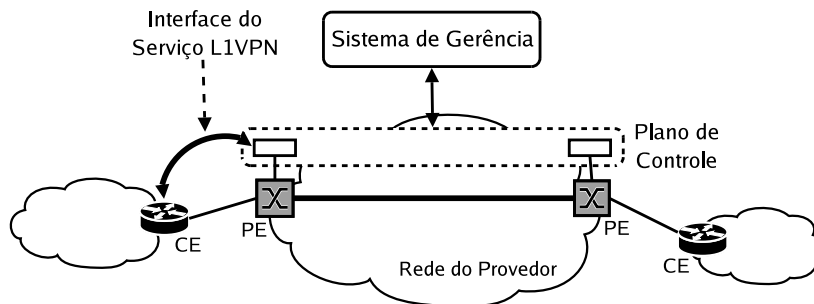


Figura 3.9: Modelo de Serviço Baseado em Sinalização.

Os elementos de borda do cliente (CEs) implementam mecanismos de sinalização, de forma que requisições de conexão podem ser enviadas diretamente ao plano de controle da rede do provedor. Neste caso, todas as operações para estabelecimento de conexões são efetuadas através de mensagens de controle e, portanto, as conexões são denominadas *switched connections*.

Uma interface baseada em um plano de controle padronizado assegura interoperabilidade e provê recursos como requisição dinâmica de conexões, notificação de falhas e mecanismos de proteção. No entanto, os protocolos e interfaces do plano de controle devem ser estendidos para suportar as funções do serviço L1VPN. Além disso, é necessário um mecanismo para distinguir as mensagens de controle enviadas dos CEs para os PEs, identificando-se a qual das diferentes VPNs elas pertencem. No contexto da arquitetura GMPLS, o modelo baseado em sinalização pode fazer uso da especificação da interface UNI (*User-Network Interface*), como definido na IETF [41]. A UNI consiste na interface entre uma rede cliente e uma rede de transporte em um cenário *overlay*.

O **modelo baseado em sinalização e roteamento** é similar ao anterior, entretanto, além de mensagens de controle de sinalização, a interface do serviço suporta troca de **informações de roteamento**. Assim, neste modelo, há adjacências de roteamento entre CE e PE. A Figura 3.10 caracteriza este modelo.

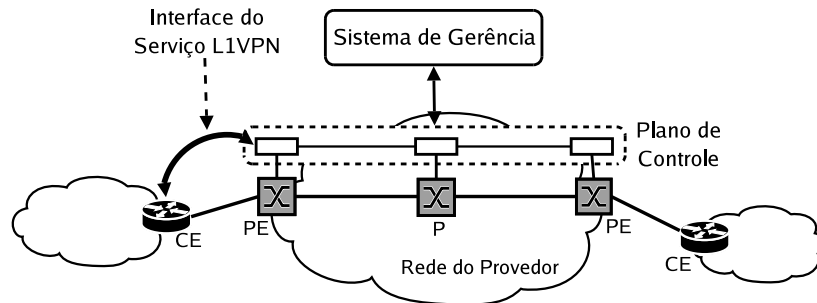


Figura 3.10: Modelo de Serviço de Sinalização e Roteamento.

Neste caso, um CE pode obter do PE, via plano de controle, informações de roteamento das redes clientes remotas que pertencem à mesma VPN, bem como informações (limitadas) de roteamento da rede do provedor. Essas informações de roteamento podem incluir informações de “alcançabilidade” (*reachability information*) ou estado de enlaces para suporte a engenharia de tráfego. Um PE deve ser capaz de manter instâncias de roteamento diferentes para cada VPN (que possui um CE conectado a esse PE).

Este modelo suporta um cenário *peer*, no caso dos elementos da rede cliente também implementarem os mecanismos do plano de controle. Além disso, as informações de roteamento dependem do modelo de alocação de recursos. Por exemplo, no caso de um serviço L1VPN com alocação dedicada, o cliente pode receber informações detalhadas da partição da rede para ele reservada.

Existe ainda algumas derivações deste modelo que dependem principalmente do nível de abstração das informações de roteamento que são fornecidas aos nós da rede cliente. Por exemplo, o provedor informa ao cliente uma topologia virtual ou toda a rede do provedor é vista pelo cliente como um único nó virtual.

Capítulo 4

Gerência Baseada em Políticas

De fato, as redes de computadores são sistemas muito complexos. Não é fácil lidar com as tarefas de monitorar e configurar os diversos componentes de *hardware* e *software* envolvidos. Além disso, fatores como ocorrência de falhas, sub-utilização de recursos, dinâmica do tráfego, aplicações com diferentes requisitos, segurança, entre outros, evidenciam a necessidade de metodologias e mecanismos que possibilitem uma **utilização eficiente dos recursos** da infra-estrutura de comunicação.

Um papel fundamental da gerência de redes é justamente oferecer essas metodologias e mecanismos. Os procedimentos de monitoração e controle compreendem vários aspectos, de maneira que a gerência de redes é dividida em cinco áreas funcionais: gerência de falhas, gerência de configuração, gerência de contabilidade (*accounting*), gerência de desempenho e gerência de segurança. Uma definição bastante abrangente para gerência de redes é apresentada em [25]:

“Gerência de redes inclui o oferecimento, a integração e a coordenação de elementos de hardware, software e humanos, para monitorar, testar, consultar, configurar, analisar, avaliar e controlar os recursos da rede, e de elementos, para satisfazer às exigências operacionais, de desempenho e de qualidade de serviço em tempo real a um custo razoável.”

Em linhas gerais, a infra-estrutura de gerência tradicional em redes IP é baseada em um modelo cliente-servidor [25]. Sua arquitetura apresenta três elementos principais: entidade gerenciadora, entidade gerenciada e protocolo de gerência de rede. A Figura 4.1 ilustra essa arquitetura. A **entidade gerenciada** compreende um elemento de rede habilitado com um agente de gerência. Esse agente é capaz de manipular (ler ou alterar) as informações de gerência sobre os componentes físicos e lógicos (como protocolos) do elemento de rede. Essas informações consistem nos parâmetros que determinam o estado e a

operação do elemento de rede. Elas são representadas conforme um modelo de dados padronizado e armazenadas em uma Base de Informações de Gerência (MIB – *Management Information Base*).

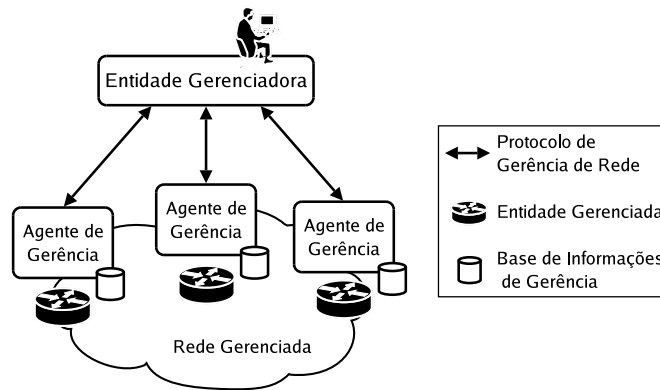


Figura 4.1: Arquitetura de gerência de redes IP convencionais.

A **entidade gerenciadora** é responsável por coletar e analisar as informações de gerência dos elementos de rede. Esse procedimento é realizado através de um **protocolo de gerência**, que suporta a comunicação entre a entidade gerenciadora e os agentes de gerência. Por meio desse protocolo, a entidade pode obter as informações de gerência dos elementos de rede, consultando os agentes correspondentes. Esse protocolo também suporta o controle da operação da rede, pois ele permite que a entidade gerenciadora configure (modifique) o estado dos elementos de rede. Para tanto, a entidade envia requisições aos respectivos agentes, que então alteram os parâmetros correspondentes. O protocolo de gerência mais utilizado em redes IP é o SNMP (*Simple Network Management Protocol*). Além das operações de consulta e alteração das informações de gerência, esse protocolo possibilita que os agentes enviem mensagens não solicitadas com o propósito de notificar a entidade gerenciadora sobre eventos ou mudanças no estado do elemento de rede.

Como enfatizado em [25], **o protocolo de gerência não gerencia a rede**. Ele apenas permite que a entidade gerenciadora monitore (consulte) ou configure (modifique) o estado dos elementos de rede. De fato, é o administrador quem realiza as decisões de gerência. Através da entidade gerenciadora ele é capaz de monitorar e analisar as informações de gerência, assim como efetivar ações de controle, reativas ou pró-ativas, para ajustar a configuração dos elementos de rede, em função de objetivos operacionais, alterações no ambiente, situações adversas, etc.

No entanto, essa abordagem demonstrou-se ineficiente para lidar com a complexidade dos processos de gerência, o que ficou evidente com a proliferação das redes de comunicação e o desenvolvimento de serviços e aplicações avançados. Um dos principais problemas

refere-se ao **alto grau de intervenção humana** no processo de reconfiguração da rede. Operações manuais são lentas e propensas a erros. A gerência de configuração envolve **altos custos** devido à heterogeneidade de equipamentos e serviços e ao baixo nível de automatização (o que exige grande número de profissionais altamente qualificados).

Além disso, o baixo grau de automatização inviabiliza que o sistema de gerência seja capaz de estabelecer uma relação entre gerência de falha e desempenho com gerência de configuração, a fim de que o sistema possa reagir dinamicamente a condições adversas. Os sistemas de gerência não suportam a especificação de objetivos administrativos e operacionais de alto-nível e nem a aplicação automática dos mesmos em procedimentos de reconfiguração da rede.

Outro problema é que essa abordagem está focada na gerência do elemento de rede. Entretanto, a diversidade de elementos e os diferentes requisitos das aplicações exigem que a gerência seja baseada em metodologias e mecanismos que considerem a rede como um todo, os serviços e os objetivos administrativos e de negócios. Neste contexto, gerência pode ser considerada em diferentes níveis, conforme ilustrado na Figura 4.2.

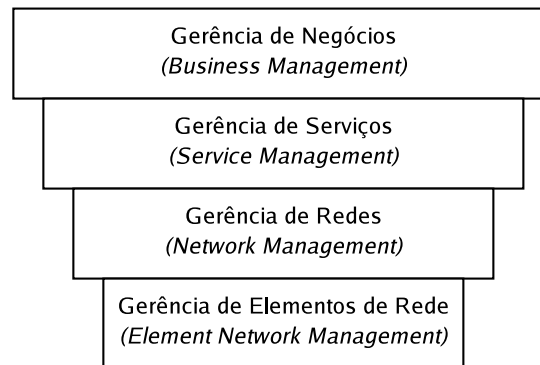


Figura 4.2: Níveis de gerência.

Esses problemas ressaltaram a necessidade de novas abordagens para a gerência de redes. Uma solução proposta, no contexto da IETF, foi a abordagem de **Gerência Baseada em Políticas** – PBM (*Policy-Based Management*) [40]. Ela foi motivada primeiramente por duas necessidades: (1) automatizar os processos de gerência de configuração de redes e serviços; (2) oferecer mecanismos para o suporte a Qualidade de Serviço (QoS) diferenciada para os diversos serviços que compartilham a infra-estrutura de rede. A abordagem PBM tem sido utilizada com sucesso para lidar com a complexidade de gerência de redes. Ela provê gerência da rede como um todo, de forma consistente e dinâmica. Como apresentado em [40], PBM pode ser definida como o uso de políticas de alto-nível para gerenciar a configuração e o comportamento de uma ou mais entidades, assim como para realizar decisões de gerência de forma dinâmica. Por sua vez, uma política é um conjunto

de regras utilizadas para gerenciar e controlar a alteração (ou não) do estado de um ou mais objetos gerenciados.

4.1 Visão Geral

Na abordagem PBM, a gerência considera também os serviços e a rede como um todo, e não somente cada elemento de rede em particular [10]. A Figura 4.3 apresenta uma visão geral da abordagem PBM. A administração de rede define um **conjunto de políticas** que controlam o acesso e a utilização dos recursos da rede, assim como a operação de protocolos e serviços. As políticas definem, por exemplo, regras de prioridade e controle de acesso a recursos e aplicações, mecanismos de escalonamento, configuração de alocação de recursos, entre outros. Ainda, elas devem ser especificadas de forma a refletir os objetivos, requisitos e prioridades pertinentes, segundo o modelo de negócio e os objetivos administrativos.

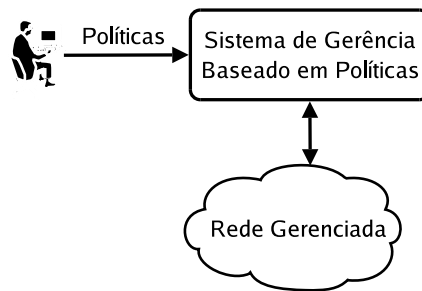


Figura 4.3: Visão geral da abordagem PBM.

Nesse contexto, o administrador delega a tarefa de realizar as decisões de gerência para o sistema. Para tanto, ele define políticas de alto nível que controlam a operação da rede. Essas políticas devem refletir os requisitos dos serviços, assim como os objetivos administrativos e corporativos. Além disso, elas definem como o sistema deve reagir a eventos e condições específicos. Dessa forma, o sistema pode responder **automaticamente** a determinadas condições e eventos, executando ações de reconfiguração apropriadas, conforme descrito pelas políticas. Portanto, um sistema de gerência baseada em políticas deve ser capaz de **traduzir** as regras de alto nível em comandos de configuração para os respectivos componentes da rede. Alguns exemplos de políticas que consideram esses aspectos são apresentados a seguir:

“Dados com origem do setor financeiro devem ser criptografados.”

“Reserve 20% da largura de banda disponível para tráfego FTP.”

“Se a aplicação é vídeo-conferência então utilizar a classe de serviço ‘Ouro’.”

A abordagem PBM se apresenta como uma metodologia simples e unificada para controlar a rede, onde a abstração através de políticas de alto nível torna transparente as particularidades de tecnologias e dispositivos heterogêneos. Como descrito em [40], essa abordagem foi proposta como meio para:

- Prover serviços diferenciados para diferentes usuários;
- Simplificar a gerência de serviços, de rede e de dispositivos;
- Diminuir o número de profissionais especializados necessários para configurar a rede;
- Definir o comportamento de uma rede ou de um sistema distribuído;
- Lidar com a crescente complexidade na operação e configuração dos equipamentos;
- Utilizar procedimentos, objetivos e requisitos de negócios (*business goals*) para dirigir a configuração da rede.

É importante ressaltar que a abordagem de gerência baseada em políticas, não necessariamente deve substituir sistemas com a arquitetura cliente-servidor, discutidos anteriormente. Ela pode ser utilizada de forma complementar no sentido de automatizar funcionalidades ou acrescentar novas. Assim, sistemas de gerência baseados em políticas podem comunicar com um sistema de gerência de redes presente na rede.

4.2 Modelos de Informação de Políticas

Como apresentado na seção anterior, em uma arquitetura de gerência, as entidades gerenciadas devem prover informações que possam ser lidas e modificadas por um sistema de gerência (a entidade gerenciadora) que então é capaz de monitorar e controlar a rede. As informações de gerência precisam estar de acordo com um modelo de informação. Portanto, o *framework* de gerência precisa prover modelos de informações de gerência. Esses modelos devem considerar o escopo das entidades gerenciadas, segundo os níveis de gerência (Figura 4.2). Além disso, uma vez que as redes são sistemas complexos e heterogêneos, o compartilhamento de informações de gerência é um desafio. Portanto, existe a necessidade de **modelos de informação padronizados**. No mesmo sentido, as políticas também devem ser especificadas de acordo com modelos próprios.

Em um ambiente de gerência baseada em políticas, modelos de informação são definidos para representar informações de políticas. Dessa forma, políticas podem ser compartilhadas entre diferentes sistemas, desde que eles considerem o mesmo modelo. Neste contexto, a IETF especificou um modelo denominado PCIM (*Policy Core Information*

Model) [35, 34]. Ele é um modelo de informação orientado a objetos que define hierarquias de classes que permitem representar informações de políticas e controle de políticas. Essas classes foram definidas com o intuito de servir como uma hierarquia base, para que outros modelos possam ser definidos a partir de sua extensão. Desta maneira, modelos de políticas para **domínios** de aplicação específicos, como QoS ou segurança, podem ser especificados como extensão do modelo PCIM, de forma que sub-classes específicas considerem aspectos específicos do domínio em questão.

O modelo deve assegurar que políticas podem ser definidas em **diferentes níveis de abstração**, segundo os níveis de gerência. Por exemplo, podemos ter políticas para gerência de serviços e outras para gerência de elementos de rede. Desta maneira, políticas podem ser definidas em um nível de abstração que é independente de tecnologias ou dispositivos específicos. Além disso, deve haver uma metodologia bem definida para mapeamento entre políticas de diferentes níveis. Um modelo de informação também deve ser independente da tecnologia para armazenar a informação. Assim o mesmo modelo de informação pode ser representado por diversos modelos de dados. Esses aspectos são importantes para garantir **interoperabilidade** no compartilhamento e reutilização de informações e políticas de gerência.

No modelo PCIM, uma política consiste em um conjunto de regras que são formadas por condições e ações. As condições são avaliadas mediante uma requisição ou quando da ocorrência de determinados eventos. Se as condições são satisfeitas, as ações são então executadas. Os elementos definidos no modelo permitem especificar condições compostas, formadas por mais de uma expressão condicional. O grupo de expressões em uma condição composta pode ser definido como conjuntivo ou disjuntivo. No primeiro caso, todas as expressões precisam verdadeiras para que a condição seja avaliada como verdadeira. No segundo, basta que uma das expressões seja verdadeira.

Além disso, as políticas podem ser agrupadas e organizadas de forma hierárquica, por meio do uso de agregação ou aninhamento (*nesting*). Prioridades também podem ser definidas para as políticas. O propósito do modelo PCIM não é definir o conteúdo, mas sim a estrutura das políticas. Neste sentido, o modelo inclui propriedades que permitem controlar o uso das políticas, por exemplo, através da definição da ordem de execução ou da prioridade das ações, ou ainda do período em que uma política é válida. A Figura 4.4 apresenta algumas classes do modelo PCIM que descrevem essa estrutura de políticas.

Além de modelos de informação, muitas linguagens foram propostas para especificação de políticas. Algumas permitem representar eventos relacionados com determinadas políticas. Um exemplo de linguagem para especificação de políticas para gerência de redes e sistemas distribuídos é a linguagem PONDER [12]. Um fator muito complexo no contexto de representação de políticas é lidar com verificação e resolução de conflitos entre políticas. Alguns desses aspectos são abordados no trabalho apresentado em [13],

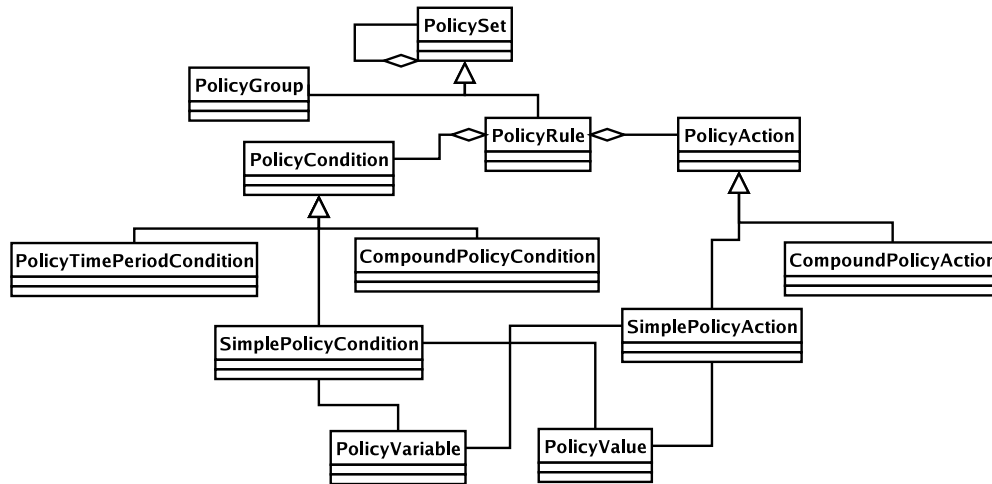


Figura 4.4: Algumas classes do modelo PCIM.

que propõe mecanismos para gerência de políticas considerando a linguagem PONDER.

4.3 Framework

Foi apresentado que em um arquitetura de gerência baseada em políticas um administrador define políticas que são aplicadas na rede a fim de controlar o comportamento do sistema como um todo. No *framework* de políticas da IETF, as políticas são regras que controlam a utilização dos recursos e serviços da rede. Essas regras definem quais ações devem ser executadas em determinadas condições. As políticas são então representadas como **expressões condicionais**, que associam um conjunto de condições a um conjunto de ações, de maneira a determinar quais ações devem ser executadas caso certas condições sejam (ou não) satisfeitas.

Um aspecto muito importante neste *framework* é a possibilidade de definição de papéis (*roles*) par as entidades gerenciadas [35]. Um papel pode ser entendido como um tipo de atributo que permite selecionar quais políticas são aplicáveis a um determinado conjunto de entidades (aquelas que “desempenham” aquele papel).

Um administrador associa a cada entidade gerenciada um ou mais papéis e então determina quais políticas estão relacionadas a cada papel. O sistema é responsável por configurar cada entidade de acordo com as políticas relacionadas com o papel ao qual esta entidade está associada. Quando o administrador altera uma política, as novas configurações são aplicadas a todas as entidades associadas aos papéis com os quais aquela política está relacionada. Isso confere maior escalabilidade e flexibilidade ao sistema.

Além disso, o conceito de papéis possibilita que a definição de um política seja dire-

cionada para uma funcionalidade e não para características ou tipos específicos de componentes da rede. Alguns exemplos de papéis são: roteador de borda, interface de rede e serviço VPN.

Em linhas gerais, a aplicação de políticas por um sistema PBM envolve monitorar o estado da rede, avaliar as condições e executar as ações correspondentes. A Figura 4.5 apresenta esse *framework*, que serve como um modelo conceitual para o projeto de sistemas de gerência baseada em políticas [40]. Como descrito pela figura, a arquitetura geral de um sistema PBM inclui três entidades principais: elemento de gerência de políticas; Ponto de Decisão de Políticas – PDP (*Policy Decision Point*); e Ponto de Aplicação de Políticas – PEP (*Policy Enforcement Point*).

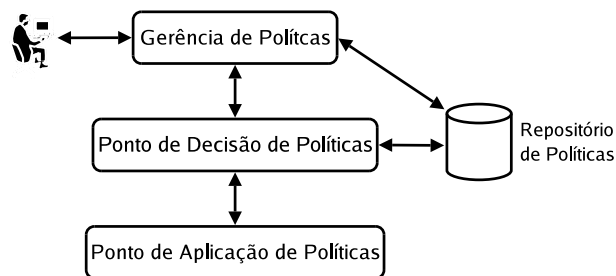


Figura 4.5: *Framework* da abordagem PBM.

Ferramenta de Gerência: este elemento é responsável pelo suporte à gerência de políticas. Ele permite ao administrador de redes criar, remover ou editar políticas, assim como provê suporte ao controle do ciclo de vida das políticas. A ferramenta de gerência de políticas considera um modelo de informações e implementa esse modelo através de uma linguagem para representação das políticas. Isso inclui mecanismos para correção sintática e semântica e para resolução de conflitos. As políticas são então armazenadas em um repositório de políticas. Esse repositório pode ser lógica ou fisicamente distribuído para armazenar políticas de diferentes domínios de aplicação.

PDP: é a entidade responsável por realizar decisões de gerência considerando as regras de políticas, o estado da rede, sua topologia e as características dos elementos da rede. Ele avalia as políticas e converte as ações em decisões de gerência para serviços e elementos de rede. Neste processo, as políticas de alto nível de abstração são traduzidas para dados de configuração específicos para as entidades gerenciadas. Isso envolve testar as condições das políticas e então gerar comandos de configuração a partir das ações correspondentes. As decisões de configuração podem ser requisitadas ou podem ser efetuadas em resposta à ocorrência de determinados eventos, como

a requisição de conexão em um elemento de borda da rede, ou a notificação de que um limiar de controle foi ultrapassado. Os comandos de configuração são gerados considerando as propriedades específicas dos elementos de rede e dos serviços. O elemento de gerência de políticas e o PDP, junto com o repositório, às vezes são considerados como um único componente denominado Servidor de Políticas (*Policy Server*).

PEP: este elemento está presente nas entidades gerenciadas. Ele é responsável por efetivamente executar as ações correspondentes às decisões definidas e enviadas por um PDP. Outra função do PEP é reportar informações de estado e características das entidades gerenciadas para o PDP.

O **protocolo COPS** (*Common Open Policy Service*) [14] foi definido pela IETF para a comunicação de informações de políticas entre PEP e PDP. Ele é um protocolo de requisição e resposta. Entre suas características pode-se destacar: segue um modelo cliente-servidor; utiliza o TCP como protocolo de transporte; é extensível, no sentido em que suporta objetos auto-identificáveis que possibilitam a troca de informações de políticas sobre clientes e domínios de aplicação específicos.

A Figura 4.6 ilustra um cenário típico, no qual um roteador de borda requisita uma decisão de política para o procedimento de controle de admissão, ao receber uma solicitação de reserva de recursos através de um protocolo de sinalização como o RSVP. Quando a solicitação chega (1), o PEP requisita (2) ao PDP uma decisão de gerência para controle de admissão. Após receber a resposta (3), e supondo que a solicitação foi aceita, a solicitação de reserva de recursos (mensagens *path*) é propagada até o destino. Mensagens de confirmação da reserva (*resv*) são enviadas no sentido contrário, até o nó de origem da solicitação.

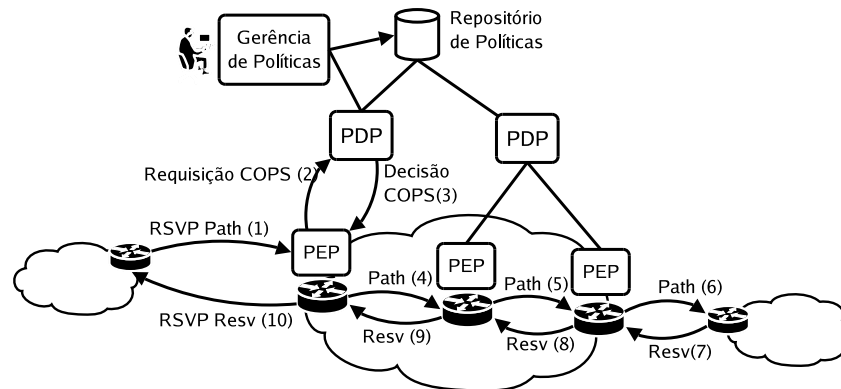


Figura 4.6: Cenário típico de uso do *framework* PBM.

O protocolo COPS suporta dois modelos de operação. No modelo *Outsourcing*, o PEP

requisita uma decisão de gerência ao PDP devido à ocorrência de um evento na entidade gerenciada. O PDP então envia uma resposta de acordo com as políticas definidas e as informações enviadas junto com a requisição (por exemplo, o papel da entidade gerenciada). Por fim, o PEP efetivamente aplica as configurações na entidade gerenciada. No modelo *Provisioning* (também denominado *Configuration*), o PDP pode, pró-ativamente, enviar decisões de gerência para o PEP, mesmo que ele não tenha sido requisitado para tal. A Figura 4.7 ilustra os modelos de operação.

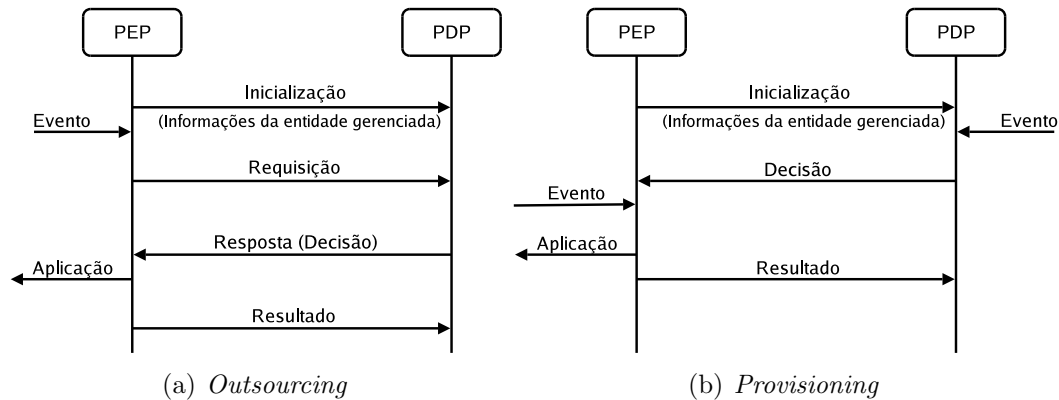


Figura 4.7: Modelos de operação do protocolo COPS.

4.4 Vantagens e Desvantagens

Grande parte dos trabalhos em gerência baseada em políticas se concentram em dois domínios de aplicação, a saber, Qualidade de Serviço e segurança. Muitos aspectos da abordagem PBM ainda são desafios para a comunidade de pesquisa e muitas questões precisam ser resolvidas para se possa explorar todo o seu potencial. Grande parte desses fatores estão relacionados com questões de **interoperabilidade**. São necessários avanços e padronizações em relação a: esquemas e modelos de informação; linguagens para representação de políticas; suporte à especificação de objetivos administrativos de alto nível; mapeamento entre níveis de abstração; resolução de conflitos; entre outros.

Outras questões estão relacionadas com a **escalabilidade** do *framework* PBM. O trabalho apresentado em [27] discute esses problemas e propõe uma nova arquitetura PBM segundo um modelo em três camadas. Além disso, não é fácil prever ou assegurar completamente o efeito de uma política sobre estado da rede como um todo. Uma política definida para controlar um serviço ou um dispositivo pode ter efeitos sobre outros componentes da rede.

Não obstante, a abordagem PBM claramente representa um avanço no sentido de

automatizar procedimentos de controle e gerência de redes e serviços. Como discutido em [10], na gerência baseada em políticas o foco está no estado desejado, diferentemente das abordagens tradicionais, nas quais o foco está nos mecanismos específicos para se obter o estado desejado. No primeiro caso, existe um **alto nível de abstração**, de maneira que o papel do administrador não é mais configurar individualmente os diferentes elementos de rede, mas sim, estabelecer políticas de alto nível para a rede como um todo. O sistema é que converte as políticas em operações específicas para os componentes da rede.

Outra vantagem apresentada naquele trabalho é que, no *framework* PBM, somente a gerência de políticas é centralizada, mas não a implementação das funções de controle e gerência, uma vez que os componentes (PDP e PEP) que realizam decisão e aplicação de ações de gerência são distribuídos. Com efeito, a abordagem PBM separa o controle das ações de gerência de sua implementação.

Outras vantagens de gerência baseada em políticas são discutidas em [40]. Comparada com práticas tradicionais de super-dimensionamento da rede, PBM representa um alternativa inteligente para gerenciar alocação de recursos e prover diferentes níveis de serviço. Essa abordagem também contribui para aprimorar aspectos de gerência de segurança, por meio de regras avançadas de controle de acesso a recursos e controle de encaminhamento de tráfego pelos roteadores. Outra vantagem são questões temporais. As políticas permitem escalonar as ações de gerência de forma flexível, assim como programar que ações ocorram em períodos (janelas de tempo) específicos.

Em suma, a **simplificação** e **automação** do processo de gerência de rede são as principais vantagens da abordagem PBM [52]. A simplificação das funções de gerência é proporcionada através de dois aspectos principais da arquitetura de políticas: centralização e abstração. O primeiro se refere ao processo de definir ações de gerência em uma entidade central, ao contrário de configurar cada elemento de rede em particular. Como discutido anteriormente, somente a especificação e gerência das políticas são centralizadas. Abstração por sua vez simplifica a definição de políticas, uma vez que assim o administrador prescinde de conhecer os detalhes de tecnologias e equipamentos. Além disso, a definição de políticas em alto-nível permite incluir os requisitos e objetivos administrativos nos processos de gerência.

A automação consiste no fato de que o administrado por si mesmo não precisa mais monitorar e configurar a rede. Ele somente define políticas que devem controlar o comportamento da rede como um todo. Então o ciclo de gerência se inicia. O servidor de políticas decide as ações de gerência apropriadas em função do estado monitorado, eventos, características dos componentes e topologia da rede. As decisões resultam em operações de configuração que são aplicadas pelos PEPs. Então novos eventos e informações sobre o estado da rede resultam em novas ações e assim o processo continua, sendo que as políticas também podem ser continuamente refinadas.

Parte II

Proposta

Capítulo 5

Contexto do Trabalho

Neste capítulo, são discutidos trabalhos em áreas relacionadas a fim de definir o contexto da presente proposta. Os trabalhos apresentados estão organizados em dois grupos. A próxima seção apresenta os trabalhos que envolvem a aplicação da abordagem de Gerência Baseada em Políticas (PBM), tanto na gerência de VPNs (de camada 2 ou 3) como no contexto de redes óticas. Um segundo grupo, apresentado na Seção 5.2, compreende os trabalhos sobre serviços L1VPN. Por fim, na terceira seção, o presente trabalho é contextualizado em relação aos trabalhos relacionados apresentados.

5.1 Gerência Baseada em Políticas

A abordagem PBM tem sido utilizada com sucesso para lidar com a complexidade na gerência de redes e serviços. Em particular, existem várias propostas de arquiteturas baseadas em políticas para gerência de VPNs de camada 2 e 3.

Uma arquitetura para gerência de serviços VPN em redes IP é proposta em [2]. A ênfase deste trabalho é em políticas para gerência de segurança. A arquitetura proposta define um servidor de políticas centralizado. Apesar de considerar a abordagem PBM, ela não é baseada no *framework* de políticas da IETF.

O trabalho apresentado em [17] também propõe uma arquitetura baseada em políticas para gerência de serviços VPN em redes IP, mas esta proposta considera o *framework* da IETF. O principal objetivo desta proposta é oferecer uma solução para a falta de funções centralizadas para gerência de IP VPNs em redes de larga escala, o que dificulta a configuração e implantação (*deployment*) desses serviços. Este trabalho apresenta uma linguagem e um modelo de informações de políticas. Ele considera a implementação de VPNs com a tecnologia IPsec (*IP Security Protocol*) e ilustra a proposta com um estudo de caso de serviços VPN entre múltiplos domínios.

Um segundo trabalho [28] também explora a abordagem PBM para gerência de serviços

VPN implementados com a tecnologia IPsec, mas o foco deste trabalho é como esta abordagem pode ser utilizada para simplificar e automatizar a gerência de mecanismos IPsec em redes IP de larga escala.

É importante também destacar diversos trabalhos de pesquisa que consideram a abordagem PBM no contexto de gerência em redes ópticas. O trabalho apresentado em [11] discute como utilizar gerência baseada em políticas para controlar os mecanismos de provisionamento de serviços em redes ópticas. A estratégia é que diretivas de alto-nível (políticas) gerenciam os mecanismos do plano de controle da rede. Outro trabalho [49] nesta mesma linha apresenta uma abordagem baseada em políticas para mecanismos que suportam o provisionamento dinâmico de conexões em redes ópticas. Estes mecanismos de controle são baseados em *Web Services* e são controlados pelos usuários (redes clientes).

A abordagem PBM também já foi utilizada como solução para diversas funções de gerência em redes ópticas, como gerência de falhas [8] e de desempenho [16]. Este último trabalho não considera o *framework* de políticas da IETF. Ele propõe estratégias para otimizar a utilização de recursos a partir de mecanismos automáticos de reconfiguração que respondem a alterações nas condições da rede. A automatização desses mecanismos depende de decisões baseadas em políticas de reconfiguração.

Outros trabalhos utilizaram a abordagem PBM em propostas de mecanismos de controle de admissão e agregação de tráfego [51, 50]. Neste caso, são definidas arquiteturas baseadas em políticas e um conjunto específico de políticas que melhoram a eficiência desses mecanismos em aspectos como utilização de recursos da rede.

5.2 Serviços L1VPN

O desenvolvimento de um plano de controle distribuído possibilitou o estabelecimento dinâmico de conexões em redes ópticas (conexões estas então denominadas *lightpaths*). Isso motivou a idéia de serviços VPN em redes ópticas, o primeiro modelo de VPN de camada 1. Os primeiros trabalhos nesta área apresentam soluções para o projeto de topologias virtuais (*virtual topology design*) que basicamente consiste em determinar o melhor conjunto de *lightpaths* para atender uma certa demanda de tráfego, considerando certas restrições (como nível de proteção das conexões).

Uma proposta para arquitetura de VPN óptica foi apresentada em [46]. O trabalho compara as abordagens distribuída e centralizada para a arquitetura a partir da definição de um modelo funcional para VPN óptica. Com foco em aspectos de roteamento, este trabalho propõe três mecanismos para implementar as funcionalidades desse serviço: extensões ao BGP (*Border Gateway Protocol*); extensões ao OSPF; e uma adaptação da abordagem de múltiplas instâncias, utilizada em VPNs de camada 3, na qual os elementos de rede mantêm diferentes instâncias de roteamento para cada VPN. Os autores

avaliam as três propostas em relação às funções suportadas e à escalabilidade. Questões relacionadas com mecanismos de sinalização ou gerência não são consideradas.

A definição de serviços L1VPN foi uma evolução do conceito de VPN óptica. Entidades internacionais têm dispensado esforços na especificação desse serviço. A ITU definiu conceitos e especificou cenários, requisitos de alto-nível e um modelo de referência para L1VPNs [18], assim como funções e arquiteturas para o suporte a serviços L1VPN [19].

No mesmo sentido, a IETF recentemente criou um grupo de trabalho cujo objetivo é especificar mecanismos para o provisionamento de serviços L1VPN em redes de transporte com um plano de controle GMPLS. As primeiras atividades deste grupo envolvem definir os requisitos e um *framework* para serviços L1VPN [42] e em analisar a aplicação de mecanismos e protocolos GMPLS no provisionamento desses serviços [43]. Os principais aspectos abordados nestas referências foram apresentados no Capítulo 3. Ainda no escopo da IETF, existem outras três importantes iniciativas em desenvolvimento:

- Em [36], é proposto um modelo para provisionamento de serviços VPN em redes GMPLS. Este trabalho descreve o serviço denominado *Generalized VPN* (GVPN). O serviço GVPN utiliza o protocolo BGP para automatizar a descoberta de informações de membros de uma VPN (*VPN reachability auto-discovery*) e utiliza protocolos GMPLS para sinalização quando do estabelecimento de conexões.
- O trabalho apresentado em [41] descreve extensões aos mecanismos de sinalização da arquitetura GMPLS para suporte a um cenário *overlay*, no qual as redes do cliente e do provedor executam diferentes instâncias do plano de controle e existe, portanto, uma interface para troca de mensagens de controle entre essas redes. As extensões propostas consideram o suporte a serviços VPN. Neste caso, as conexões VPN solicitadas pelos nós de borda da rede cliente podem ser estabelecidas através de um mecanismo de hierarquia de conexões [22, 3]. A conexão da VPN é “tunelada” através de uma conexão na rede do provedor, que é previamente estabelecida ou é criada mediante uma nova requisição de conexão para uma VPN.
- São considerados dois modos de operação do serviço L1VPN. A descrição inicial do modo básico (*L1VPN Basic Mode*) é apresentada em [15]. Neste modo, a unidade básica do serviço é uma conexão GMPLS, ou seja, um LSP. Este modo é equivalente ao serviço baseado em sinalização, ou seja, existe troca de mensagens de controle entre CE e PE. No entanto, a comunicação envolve apenas mensagens de sinalização para estabelecimento de conexões. Não existe distribuição de informações de roteamento. O trabalho mencionado se propõe a descrever o modelo operacional desse modelo do serviço, incluindo como os membros da VPN são identificados; como esses identificadores são compartilhados entre os elementos de rede; e procedimentos

de sinalização para estabelecimento e recuperação de conexões da VPN. Ainda não há propostas para o segundo modo, denominado *L1VPN Enhanced Mode*.

Outras contribuições demonstram o importante papel da arquitetura GMPLS no suporte a serviços L1VPN. Por exemplo, o trabalho apresentado em [44] discute como os mecanismos da arquitetura GMPLS podem ser utilizados para implementar as funcionalidades do serviço L1VPN. Este trabalho primeiro descreve a motivação e os conceitos básicos do serviço L1VPN e então analisa a aplicação de mecanismos GMPLS no suporte a serviços L1VPN, considerando principalmente aspectos relacionados com endereçamento, descoberta automática de membros e sinalização de conexões. Além disso, são identificadas áreas de trabalho que exigem estudo adicional, como gerência de serviços L1VPN.

Além de requisitos e modelo funcional, são apresentados três tipos de arquiteturas para serviços L1VPN em [45]. Na arquitetura centralizada, as funções de controle são implementadas em uma entidade de gerência centralizada. Na arquitetura distribuída, essas funções são distribuídas nos nós da rede (CEs, PEs e Ps). A arquitetura híbrida é caracterizada por uma combinação das duas anteriores, onde algumas funções são centralizadas e outras distribuídas. Em geral, as funções específicas para L1VPNs, como gerência de informações de membros de VPNs, são centralizadas, enquanto que funções como controle de conexões e roteamento são distribuídas.

Uma arquitetura para serviços L1VPN é proposta em [47]. Este trabalho considera o modelo baseado em gerência como o mais adequado para os passos iniciais no desenvolvimento de serviços L1VPN, uma vez que os modelos baseados em sinalização e roteamento exigem extensões aos protocolos do plano de controle. Então são avaliadas as abordagens centralizada e híbrida para definição da arquitetura do serviço L1VPN baseado em gerência. Por fim o trabalho propõe uma arquitetura híbrida com funções de controle para serviços L1VPN, devido às vantagens da arquitetura híbrida sobre a centralizada.

5.3 Contexto do Presente Trabalho

O presente trabalho não pretende trazer contribuições para a abordagem PBM em si. No entanto, uma das contribuições desse trabalho é propor uma nova e importante aplicação para essa abordagem. Como apresentado na Seção 5.1, a abordagem PBM já foi utilizada com sucesso para simplificar e automatizar a gerência de serviços VPN de camada 2 e 3 e de serviços em redes óticas. No presente trabalho, a proposta é utilizar essa abordagem para atender os requisitos de gerência de serviços VPN de camada 1.

O conceito de serviços L1VPN é relativamente novo [18]. Como apresentado na Seção 5.2, a maioria dos trabalhos aborda os problemas relacionados com mecanismos e protocolos do plano de controle. Diferentemente, o presente trabalho envolve questões relacio-

nadas com a gerência. Neste trabalho é proposta uma arquitetura baseada em políticas para a gerência de serviços L1VPN.

A arquitetura aqui proposta considera o modelo funcional e a abordagem de arquitetura híbrida apresentados em [45]. Uma proposta relacionada é a arquitetura de serviço L1VPN apresentada em [47], cujo foco são aspectos de roteamento e controle de conexões. De fato, estas propostas não são concorrentes, mas complementares. A principal contribuição da arquitetura aqui proposta é o enfoque na gerência de serviços L1VPN, principalmente na gerência de configuração, e o uso da abordagem de Gerência Baseada em Políticas como solução para prover, de forma independente, algum nível de controle e gerência para os clientes de serviços L1VPN.

Capítulo 6

Framework de Políticas

A metodologia considerada nesta proposta foi utilizar a abordagem de Gerência Baseada em Políticas a fim de garantir que cada cliente possa exercer algum nível de controle e gerência sobre sua VPN de forma independente. De fato, um requisito do serviço L1VPN é que o cliente seja capaz de especificar políticas para controlar a operação de seu serviço. Portanto, a abordagem PBM emerge como uma solução viável para atender a esses requisitos.

Neste contexto, três questões são abordadas a seguir. Primeiro, é proposto um conjunto de classes de políticas para gerência de serviços L1VPN. Segundo, é discutido como o *framework* PBM da IETF pode ser utilizado na gerência desses serviços. Por fim, é proposta uma alternativa para a representação de políticas para serviços L1VPN, através da definição de um modelo de informações e do uso de tecnologias XML.

6.1 Classes de Políticas

A seguir são propostas três classes para políticas de gerência de serviços L1VPN: políticas de configuração, admissão de controle e roteamento. O conjunto de classes proposto não pretende ser exaustivo, uma vez que outros tipos de políticas podem ser definidos. É importante notar que algumas políticas envolvem aspectos que podem ser considerados em mais de uma classe. Neste caso, a política deve ser classificada de acordo com o aspecto mais relevante. O objetivo aqui é organizar e ilustrar alguns dos principais aspectos da gerência de serviços L1VPN que podem ser controlados através de políticas.

6.1.1 Políticas de Configuração

As Políticas de Configuração são utilizadas para definir parâmetros de configuração que controlam a operação de serviços L1VPN, principalmente no que diz respeito ao estabe-

lecimento de conexões da VPN e controle de mecanismos de QoS. Essas políticas podem ser definidas de forma a refletir as especificações de um contrato de nível de serviço (SLA) entre cliente e provedor. A seguir são discutidos alguns aspectos de serviços L1VPN que podem ser controlados por meio de políticas de configuração:

- A infra-estrutura de rede do provedor pode suportar ambos os modelos de alocação de recursos: compartilhado e dedicado. Políticas de configuração podem ser definidas para: determinar o modelo de alocação utilizada para cada VPN; configurar os mecanismos de alocação de recursos; ou ainda, definir regras para reconfigurar o modelo de alocação utilizado para uma determinada VPN, em função de alterações nas condições da rede. Exemplos:

Para a VPN 100, utilizar modelo dedicado.

Se o modelo é compartilhado, alocar comprimentos de onda nos enlaces com maior número de comprimentos de onda disponíveis.

Se taxa de bloqueio da VPN 102 superar 5%, utilize o modelo dedicado.

- O provedor pode suportar diversos mecanismos de recuperação (*recovery schemes*). As políticas de configuração podem ser especializadas para tratamento de falhas, assim como em propostas de gerência de falhas em redes ópticas que utilizam a abordagem PBM [8]. Essas políticas podem definir e configurar o mecanismo de recuperação utilizado em cada VPN, assim como diferenciar os mecanismos utilizados para as conexões requisitadas por um membro específico de uma VPN. Exemplos:

Para VPN 100, utilizar restauração como mecanismo de recuperação.

Se membro da VPN é do departamento financeiro, utilizar como mecanismo de recuperação proteção com nível “1 + 1”.

- Se o provedor suporta diferentes classes de serviço L1VPN, políticas de configuração podem ser utilizadas para definir a classe para cada VPN. Essas políticas permitem inclusive a configuração dos mecanismos implementados para atender aos requisitos de qualidade de cada classe de serviço, como definição de prioridades, configuração de alocação de recursos, entre outros parâmetros. Exemplo:

Utilizar classe de serviço “bronze” para VPN 105.

Se a VPN é da classe “ouro”, o seu modelo de alocação deve ser dedicado e para ela deve ser reservado 4% da capacidade em cada enlace da rede.

6.1.2 Políticas de Controle de Admissão

No contexto de serviços L1VPN, um mecanismo de controle de admissão de conexões deve considerar também informações sobre os membros das VPNs, além de informações sobre a disponibilidade de recursos. As Políticas de Controle de Admissão podem ser utilizadas com o objetivo de definir regras ainda mais elaboradas para tratar as requisições de conexão dos clientes. Alguns desses aspectos são descritos a seguir:

- Este tipo de política permite controlar a utilização de recursos pelos clientes. Por exemplo, elas podem ser utilizadas para limitar o número de conexões por VPN ou por membro de uma VPN. Exemplos:

O número de conexões ativas para VPN 105 não pode exceder 30.

Membros do departamento administrativo da VPN 107 não têm limite de número de conexões.

Para a VPN 130, limite a 60 Gbps a largura de banda total alocada.

- Uma das características básicas de serviços L1VPN é a restrição de conectividade, ou seja, conexões são permitidas apenas entre membros da mesma VPN. Políticas de Controle de Admissão possibilitam definir restrições de conectividade adicionais mais elaboradas, incluindo restrições entre membros da mesma VPN.

Para a VPN 100, membros do departamento de engenharia de produção não podem estabelecer conexões com o departamento financeiro.

O único destino permitido para conexões na VPN 100 é o membro 1000 (força topologia em estrela).

- Os elementos de borda na rede do provedor possivelmente executam mecanismos para agregação dos fluxos de tráfego das redes clientes e instalação desses fluxos em conexões previamente estabelecidas na rede do provedor. Políticas de Controle de Admissão podem ser utilizadas para controlar esses mecanismos ou para otimizar a seleção de uma conexão onde o fluxo será acomodado quando existem muitas conexões adequadas disponíveis. Estas questões são exploradas no trabalho apresentado em [50], no qual foi proposta uma arquitetura baseada em políticas para agregação dinâmica de tráfego em redes ópticas. Exemplo:

Acomode LSPs de alta prioridade em lightpaths que não estão vazios, mas que só tenham LSPs de alta prioridade.¹

¹Neste caso, LSPs correspondem a LSPs em redes de comutação de pacotes e *lightpaths*, a LSPs em redes de comutação por comprimento de onda, como seria o caso de uma rede de transporte óptica com redes clientes IP/MPLS.

6.1.3 Políticas de Roteamento

Uma vez que as instâncias de roteamento são, de alguma forma, separadas por VPN, *Políticas de Roteamento* podem ser utilizadas para configurar diferentes procedimentos de roteamento para cada VPN. Por exemplo, para definir métricas e restrições para o cálculo de rotas. Além disso, essas políticas permitem configurar aspectos envolvidos na gerência de recursos que influenciam os mecanismos de roteamento. A seguir, são discutidos algumas dessas questões e os principais casos onde políticas de roteamento são adequadas:

- Políticas de Roteamento são úteis para controlar os mecanismos de roteamento, como na definição de algoritmos de cálculo de rotas, pesos de enlaces, hierarquia, parâmetros de protocolos de roteamento, entre outros. Exemplos:

Para a VPN 110, considere o peso dos enlaces igual a 1 (encontrar caminho com menor número de saltos).

Para VPNs da classe “prata”, considere o peso de um enlace como o inverso do número de lambdas alocados (prioriza enlaces mais ocupados).

- Esta classe de políticas pode ser utilizada no suporte a Roteamento Baseado em Restrições (*Constraint-Based Routing*), principalmente em mecanismos de engenharia de tráfego [3]. Nesta abordagem, o algoritmo de cálculo de rotas considera um conjunto de restrições relacionadas com aspectos administrativos e com requisitos de qualidade de serviço. As políticas de roteamento permitem definir essas restrições que condicionam o cálculo de rotas para as conexões requisitadas (*policy-based routing*). Essas políticas podem servir para: restringir por quais domínios, ou nós, uma rota pode passar; estabelecer critérios para mecanismos de balanceamento de carga; definir valores para parâmetros de QoS; entre outros. Além disso, mecanismos de recuperação de falhas podem utilizar essas políticas como restrições para o cálculo de caminhos disjuntos. Exemplos:

Conexões da VPN 100 não podem passar pelos nós 10, 15 e 20.

Para cálculo de rotas da VPN 110, considere apenas enlaces com mais de 30% da capacidade (banda) disponível.

- Políticas podem ser utilizadas na gerência de recursos no suporte aos modelos de alocação compartilhado e dedicado. Por questões de desempenho e escalabilidade, isso ocorre principalmente quando o cálculo de rotas já é realizado por uma entidade centralizada. Neste contexto, Políticas de Roteamento podem ser utilizadas para controlar quais, e de que forma, os recursos podem ser utilizados no cálculo de rotas. Exemplos:

Enlaces com capacidade igual ou superior a 80 Gbps não podem ser dedicados.

Os nós 50 e 60 são compartilhados somente entre as VPNs 100, 130 e 140.

- Pode ser o caso de existirem várias rotas (pré-calculadas) que satisfazem uma requisição de conexão entre um par origem e destino. Políticas de Roteamento podem ser utilizadas em procedimentos que otimizam a seleção de uma rota entre as disponíveis, por exemplo, com o objetivo de melhorar a taxa de utilização de recursos. Exemplos:

Para a VPN 200, priorize rotas que passem por elementos que suportam recuperação no nível de enlace.

Priorize rotas com menor valor acumulado para o nível de degradação do sinal óptico.

6.2 Framework

Nesta seção, é discutido como o *framework* de políticas da IETF pode ser utilizado no contexto da gerência de serviços L1VPN. O esquema está ilustrado na Figura 6.1, onde a direção do fluxo das informações é baseada no ciclo de vida das políticas. Este ciclo envolve basicamente três etapas: criar as políticas, gerar decisões de gerência e executar as ações correspondentes nas entidades gerenciadas.

Este esquema é detalhado no decorrer desta seção, a partir das funções das três entidades básicas do *framework* PBM: elemento de gerência de políticas, ponto de decisão de políticas (PDP) e ponto de aplicação de políticas (PEP).

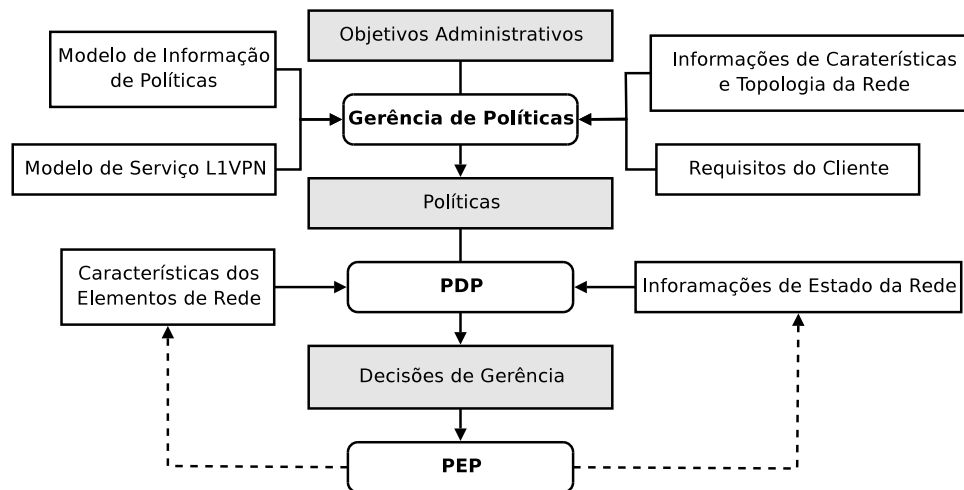


Figura 6.1: *Framework* baseado em políticas para gerência de Serviços L1VPN.

6.2.1 Elemento de Gerência de Políticas

A principal função do elemento de gerência de políticas é oferecer suporte para a especificação e controle das políticas. É através desta entidade que o administrador da rede interage com o sistema de gerência.

As políticas devem ser elaboradas segundo os objetivos administrativos e o modelo de negócios do provedor. De fato, elas devem compreender um conjunto de regras de alto nível que reflete a estratégia de gerência do provedor. Assim, as políticas podem controlar prioridades, fatores de segurança, alocação de recursos, entre outros, de acordo com os objetivos da administração da rede. Além disso, a especificação das políticas deve levar em conta os seguintes fatores:

- Modelo de serviço L1VPN: as políticas devem ser consistentes com o modelo de serviço L1VPN suportado pela rede do provedor. Algumas funcionalidades existem apenas em modelos específicos, de maneira que essas particularidades devem ser observadas quando da elaboração das políticas, principalmente no caso de gerência de configuração. Por exemplo, devem ser considerados os modelos de alocação de recursos e os tipos de interface de serviço que a rede de transporte suporta.
- Características da infra-estrutura da rede do provedor: de forma análoga, a definição das políticas deve considerar as tecnologias, topologia e outras características específicas da rede de transporte do provedor que possam influenciar na configuração ou desempenho dos serviços L1VPN.
- Requisitos do cliente: as políticas devem ser definidas de forma a atender aos objetivos e requisitos dos clientes, por exemplo, de acordo com um Contrato de Nível de Serviço – SLA (*Service Level Agreement*). De fato, as políticas precisam ser definidas para controlar as decisões de gerência dos serviços L1VPN de acordo com os objetivos e requisitos do provedor e dos clientes.

Outra questão é que as políticas sejam especificadas de acordo com um modelo de informações bem definido. Isso é importante principalmente para assegurar interoperabilidade e para flexibilizar o compartilhamento e o mapeamento das políticas. A especificação de um Modelo de Informações de Políticas é uma tarefa crítica, pois o modelo determina quais políticas podem ser elaboradas, assim como a estrutura dessas políticas. Os elementos definidos no modelo devem refletir o domínio de aplicação das políticas, neste caso, gerência de serviços L1VPN.

Além disso, elemento de gerência de políticas deve incluir suporte para a edição das políticas, incluindo funções para verificação de sintaxe, validação e resolução de conflitos. Outras funcionalidades importantes envolvem controlar o ciclo de vida das políticas, permitindo ativar, desativar ou alterar políticas existentes.

6.2.2 PDP e PEP

O resultado da etapa de especificação é um conjunto de políticas de alto nível para gerência dos serviços L1VPN. Essas políticas são então utilizadas pelos outros elementos definidos no *framework* para controlar a configuração e a operação dos serviços L1VPN.

O Ponto de Decisão de Políticas (PDP) é responsável por avaliar as políticas e gerar decisões de gerência para os elementos de rede. É importante ressaltar que o procedimento de avaliação das políticas deve considerar o estado da rede e as características dos elementos de rede. Neste processo, as políticas de alto nível de abstração são traduzidas para dados de configuração específicos para as entidades gerenciadas. As decisões de gerência podem ser requisitadas ou podem ser efetuadas em resposta à ocorrência de determinados eventos, como uma requisição de conexão.

Por fim, o Ponto de Aplicação de Políticas (PEP) é responsável por efetivamente executar nas entidades gerenciadas as ações de controle e configuração correspondentes às decisões enviadas por um PDP. Outra função do PEP é reportar informações de estado e características das entidades gerenciadas para o PDP. Estas entidades podem ser lógicas (*software*), como protocolos ou algoritmos de alocação de recursos, ou físicas, como interfaces de rede, comutadores, etc. Assim, pode ser necessário implementar um PEP em cada entidade da rede que depende das decisões de políticas, por exemplo, nos comutadores ou em um mecanismo específico do sistema de gerência, como um mecanismo de controle de admissão.

O processo de decisão executado pelo PDP é logicamente centralizado. Isso significa que podem existir diversos PDP, dependendo da escala da rede. No entanto, o controle é centralizado sob a perspectiva de que cada PDP deve acessar um único conjunto de políticas. O processo de aplicação das políticas é distribuído pela rede através da implementação das funcionalidades do PEP nas entidades gerenciadas.

6.3 Representação de Políticas

Foi discutido a importância de um modelo de informações de políticas. Uma alternativa é estender o PCIM, o modelo base de políticas da IETF. Neste caso, a extensão significa adicionar ao modelo base elementos específicos para definição de políticas para serviços L1VPN. Com efeito, o modelo base é especializado para uma aplicação específica. Assim, a estrutura das políticas e a maneira como elas são organizadas continua como no PCIM, ou seja, as políticas são definidas em termos de regras condicionais. Além disso, são reutilizados os mecanismos do PCIM que permitem definir hierarquias e grupos de políticas, prioridades, condições compostas, período em que uma política é válida, entre outros (como introduzido na Seção 4.2).

Neste contexto foi definido um modelo simplificado de informações de políticas para serviços L1VPN, conforme apresentado na Figura 6.2. Este modelo contempla elementos que permitem definir algumas políticas de configuração, admissão de controle e roteamento, conforme as classes propostas na seção anterior.

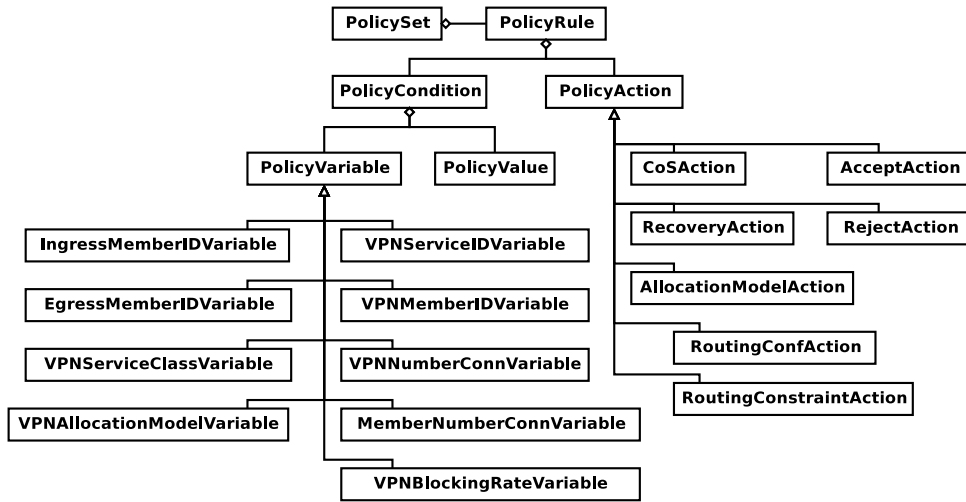


Figura 6.2: Modelo de informações de políticas para serviços L1VPN.

As políticas consistem em um conjunto de regras. Cada regra representa uma expressão condicional formada por ações (*PolicyAction*) e condições (*PolicyCondition*). As condições das regras (*PolicyRule*) são formadas associando-se uma variável (*PolicyVariable*) a um valor (*PolicyValue*). O modelo define as seguintes ações:

- Regras de configuração: ações para configurar classes de serviço (*CoSAction*), mecanismos de recuperação (*RecoveryAction*) e modelos de alocação de recursos (*AllocationModelAction*);
- Regras de controle de admissão: ações para aceitar (*AcceptAction*) ou rejeitar (*RejectAction*) conexões;
- Regras de roteamento: ações para configurar mecanismos de roteamento (*RoutingConfAction*) ou definir restrições para cálculo de rotas (*RoutingConstraintAction*).

Cada uma dessas classes que especializam a classe *PolicyAction* precisa definir os atributos necessários para especificar a ação correspondente. Por exemplo, a classe *AllocationModelAction* deve conter atributos que possibilitem definir a configuração dos mecanismos de alocação para cada serviço L1VPN.

Outro fator importante é como representar as políticas. Uma tendência na definição de arquiteturas de gerência é o uso de tecnologias XML (*Extensible Markup Language*)

[38], principalmente para gerência de configuração e para representação de informações de gerência. As vantagens de XML incluem: flexibilidade e portabilidade; ser inteligível para seres humanos; suporte a extensões e facilidade para adição de novos tipos de informação; independência da forma como os dados são apresentados; esquemas para verificação de sintaxe e consistência dos dados; mecanismos de mapeamento para outras formas de representação; suporte a informações da semântica dos dados.

A tecnologia XML é uma solução adequada para a representação de políticas de alto nível, principalmente, por fatores relacionados a: interoperabilidade; flexibilidade; suporte para representação da semântica dos dados; além da ampla aceitação da linguagem XML e da disponibilidade de diversas ferramentas. Neste contexto, a estrutura dos documentos XML deve ser especificada de acordo com o modelo de informações de políticas adotado. Para isso podem ser utilizados esquemas XML (*XML schema*) ou definições de tipos de documentos – DTD (*Document Type Definition*).

Além disso, mapeamentos entre modelos ou formas de apresentação das políticas XML podem ser obtidos com mecanismos XSLT (*Extensible Stylesheet Language Transformations*) (mecanismo para mapeamento entre documentos XML com diferentes estruturas ou para mapeamento de um documento XML em um documento em outro formato, como texto ou HTML). De fato, muitos trabalhos em gerência baseada em políticas utilizam as tecnologias XML para representação de políticas [9, 1, 33],

Considerando o modelo de políticas proposto anteriormente, a Figura 6.3 apresenta um exemplo da representação XML de uma política para serviços L1VPN. Essa política exemplo define quatro regras, duas de configuração, uma de admissão de controle e uma de roteamento:

1. A primeira regra define ações de configuração para o serviço L1VPN cujo identificador é “100”, como está estabelecido na condição da regra (linhas 6 a 9). Os valores dos atributos especificam parâmetros para configuração do modelo de alocação de recursos, da classe de serviço e do mecanismo de recuperação (linhas 11 a 13, respectivamente). Esta regra pode ser interpretada como: “para o serviço L1VPN 100, o modelo de alocação deve ser compartilhado, a classe de serviço deve ser ‘gold’, segundo o modelo ‘basic’, e o mecanismo de recuperação, proteção com nível ‘1:1’.”
2. A ação de configuração definida na segunda regra se aplica ao mesmo serviço, mas para um membro específico, cujo identificador é “CPI1-PPI1” (como pode ser observado na condição da regra, linhas 17 a 22). A ação determina qual mecanismo de recuperação deve ser utilizado para este membro da VPN, no caso, proteção com nível “1+1”. Note que existe um conflito com a regra anterior, que define um mecanismo com nível diferente. Mecanismos de prioridades, herdados do modelo PCIM, podem ser utilizados para resolver isso quando da requisição de uma decisão

de gerência sobre qual mecanismo deve ser utilizado. Os atributos herdados, que permitem definir prioridades, não estão representados na figura. Essa regra pode ser interpretada como: “para o membro ‘CPI1-PPI1’ da VPN ‘100’, utilize proteção de nível ‘1+1’ como mecanismo de recuperação.”

3. A terceira regra da política é uma regra de controle de admissão que define um exemplo de restrição de conectividade. Neste caso, a regra foi utilizada para impedir que uma conexão seja estabelecida entre os membros “CPI2-PPI2” e “CPI3-PPI3” da VPN “100”, como especificado na condição (linhas 28 a 35). A ação determina que uma requisição de conexão entre esses membros deve ser rejeitada (linha 37). Uma interpretação para essa regra seria: “Para a VPN 100, se o membro de ingresso é ‘CPI2-PPI2’ e o de egresso é ‘CPI3-PPI3’, então a conexão deve ser rejeitada.
4. A última regra define ações para configuração de roteamento e definição de restrições no cálculo de rotas. A condição da regra determina qual serviço L1VPN dever ser considerado, no caso, aquele com identificador igual a “200” (linhas 41 a 44). Duas ações são especificadas: a primeira determina o algoritmo para cálculo de rota e a métrica para peso dos enlaces (linha 46); a segunda determina restrições para o cálculo de rotas (linha 47). Neste caso, as restrições especificam que uma rota para a VPN 200 não pode passar por um dos elementos de rede “1000”, “1010” ou “1020” e que os elementos na rota devem suportar pelos menos o mecanismo de proteção com nível “1:1” .

```

1 <?xml version='1.0'?>
2 <!DOCTYPE Policy SYSTEM "11vpnPolicy.dtd">
3 <Policy id="001">
4   <PolicySet>
5     <PolicyRule type="configuration">
6       <SimplePolicyCondition>
7         <VPNServiceIDVariable/>
8         <PolicyValue>100</PolicyValue>
9       </SimplePolicyCondition>
10      <CompoundPolicyAction>
11        <ResourceAllocationAction allocationModel="shared"/>
12        <CoSAction model="basic" class="gold"/>
13        <RecoveryAction recoveryScheme="protection" level="1:1" />
14      </CompoundPolicyAction>
15    </PolicyRule>
16    <PolicyRule type="configuration">
17      <CompoundPolicyCondition>
18        <VPNServiceIDVariable/>
19        <PolicyValue>100</PolicyValue>
20        <VPNMemberIDVariable/>
21        <PolicyValue>CPI1-PPI1</PolicyValue>
22      </CompoundPolicyCondition>
23      <SimplePolicyAction>
24        <RecoveryAction recoveryScheme="protection" level="1+1" />
25      </SimplePolicyAction>
26    </PolicyRule>
27    <PolicyRule type="admissionControl">
28      <CompoundPolicyCondition>
29        <VPNServiceIDVariable/>
30        <PolicyValue>100</PolicyValue>
31        <IngressMemberIDVariable/>
32        <PolicyValue>CPI2-PPI2</PolicyValue>
33        <EgressMemberIDVariable/>
34        <PolicyValue>CPI3-PPI3</PolicyValue>
35      </CompoundPolicyCondition>
36      <SimplePolicyAction>
37        <RejectAction/>
38      </SimplePolicyAction>
39    </PolicyRule>
40    <PolicyRule type="routing">
41      <SimplePolicyCondition>
42        <VPNServiceIDVariable/>
43        <PolicyValue>200</PolicyValue>
44      </SimplePolicyCondition>
45      <CompoundPolicyAction>
46        <RoutingConfAction pathCompAlgorithm="dijkstra" linkWeight="1/w"/>
47        <RoutingConstraintAction notAllowedNode="[1000,1010,1020]"
48          minRecoveryScheme="protection" minRecoveryLevel="1:1"/>
49      </CompoundPolicyAction>
50    </PolicyRule>
51  </PolicySet>
52 </Policy>

```

Figura 6.3: Exemplo de política em XML.

Capítulo 7

Arquitetura Proposta

Neste capítulo, é apresentada a proposta de uma arquitetura baseada em políticas para a gerência de serviços L1VPN. Primeiro, será descrito o modelo funcional da arquitetura e apresentado um exemplo de como os módulos da arquitetura interagem entre si a fim de esclarecer melhor o seu funcionamento.

Na Seção 7.2, são discutidos dois cenários de aplicação da arquitetura. Estes cenários caracterizam as diferentes funções que podem ser desempenhadas pelas redes clientes e pelo provedor. O primeiro cenário corresponde a uma abordagem centralizada, na qual a rede cliente acessa as funcionalidades do serviço L1VPN através de uma interface de gerência. Por outro lado, o segundo cenário descreve uma abordagem distribuída, na qual as funcionalidades do serviço L1VPN são providas através de protocolos do plano de controle.

O projeto da arquitetura assume que a rede de transporte possui um plano de controle que provê controle dinâmico de conexões e mecanismo de roteamento com descoberta automática de informações de topologia. Em particular, para uma explicação mais concreta, será considerada uma rede de transporte óptica com um plano de controle GMPLS.

7.1 Modelo Funcional

A arquitetura proposta é apresentada na Figura 7.1. Suas funcionalidades foram definidas em função dos requisitos de serviços L1VPN. A arquitetura possui duas interfaces que isolam os módulos principais, com o objetivo de prover maior grau de flexibilidade e de interoperabilidade. Essas interfaces são descritas a seguir:

Interface de Acesso: Esta é a interface através da qual provedor ou clientes acessam as funções de gerência dos serviços L1VPN. Sua responsabilidade é processar as requisições desses usuários. Eles podem acessar esta interface para requisitar co-

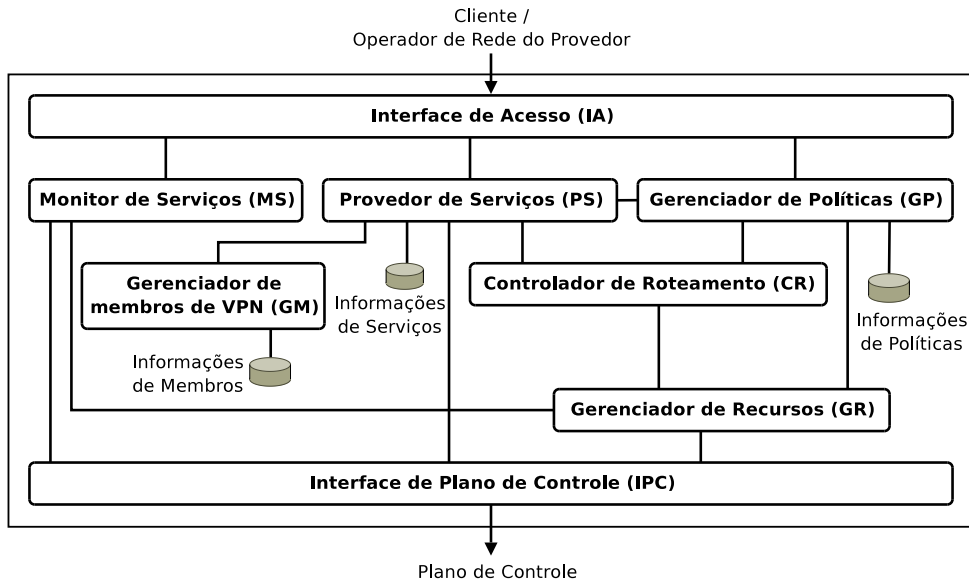


Figura 7.1: Arquitetura para gerência de serviços L1VPN.

nexões, definir políticas ou receber informações sobre o serviço L1VPN correspondente. Além disso, esta interface deve suportar mecanismos de controle de acesso e autenticação.

Interface de Plano de Controle: O sistema de gerência acessa as funcionalidades do plano de controle a partir desta interface, a fim de controlar o estabelecimento de conexões, obter informações de roteamento, de disponibilidade de recursos e notificações de ocorrência de falhas. Portanto, essa interface depende da arquitetura de plano de controle implementada na rede do provedor. Alterações no plano de controle do provedor implicam alterações nesta interface, mas não necessariamente nos módulos internos da arquitetura.

Os administradores interagem com o sistema através da Interface de Acesso. Em função das requisições processadas, esta interface aciona algum dos três módulos descritos a seguir:

Monitor de Serviços: Este módulo é responsável por gerenciar informações sobre desempenho e ocorrência de falhas. Ele permite que cada cliente seja capaz de monitorar as informações relacionadas a seu serviço. Essas informações são obtidas diretamente do plano de controle ou a partir da análise e correlação de dados sobre os recursos da rede, mantidos pelo próprio sistema de gerência (no Gerenciador de Recursos).

Provedor de Serviços (PS): Este é o módulo principal da arquitetura. Ele é responsável por “instanciar” e configurar serviços L1VPN. É também sua responsabilidade o controle de admissão sobre requisições de conexão. Além disso, este módulo se encarrega das funções de controle de conexões, como criar, recuperar ou remover conexões entre membros da VPN. Para isso ele invoca os mecanismos apropriados do plano de controle. Este módulo mantém informações sobre as conexões que formam uma determinada VPN.

Gerenciador de Políticas (GP): Este módulo deve incluir as funcionalidades de uma ferramenta de gerência de políticas e de um PDP, como descrito no capítulo anterior. Assim, ele permite aos clientes e provedor adicionar, remover ou alterar as políticas para gerência dos serviços. Outra função sua é responder a requisições de decisões de gerência de outros módulos a partir das políticas especificadas, que ficam armazenadas em um repositório de políticas.

Os demais módulos da arquitetura são descritos a seguir:

Gerenciador de Membros (GM): a função deste módulo é gerenciar as informações sobre quais membros pertencem a cada VPN. Essas informações podem ser fornecidas ao sistema de forma estática (por configuração), ou, no caso de uma arquitetura distribuída, essas informações podem ser compartilhadas de forma automática (*VPN Membership Auto-Discovery*).

Gerenciador de Recursos (GR): é responsável por gerenciar informações sobre o estado da rede, mantendo informações sobre a disponibilidade de recursos. Este módulo deve oferecer suporte aos dois modelos de alocação de recursos, compartilhado e dedicado.

Controlador de Roteamento (CR): é o módulo que calcula rotas para o estabelecimento de conexões. O cálculo de rotas é realizado a partir das informações fornecidas pelo gerenciador de recursos e das informações de topologia obtidas do mecanismo de roteamento do plano de controle.

Na Tabela 7.1 apresenta-se um resumo das funcionalidades da arquitetura proposta. Essas funcionalidades podem ser implementadas como um sistema independente ou integradas ao sistema de gerência do provedor. De fato, algumas destas funcionalidades, como cálculo de rotas, não são específicas de serviços L1VPN e talvez já estejam implementadas no sistema de gerência do provedor. Neste caso, eventuais modificações ou extensões podem ser necessárias para que estas funções suportem características específicas dos serviços L1VPN.

Tabela 7.1: Principais funcionalidades da arquitetura.

<i>Módulo</i>	<i>Funções</i>
Interface de Acesso	Processar de requisições; autorização e autenticação.
Monitor de Serviços	Fornecer informação sobre desempenho e falhas.
Provedor de Serviços	Configurar serviços L1VPN; Executar controle de admissão e de conexões.
Gerenciador de Políticas	Suporte a modelagem de políticas; Processar requisições de decisões de gerência.
Gerenciador de Membros	Adicionar ou remover membros. Verificar pertinência.
Controlador de Roteamento	Calcular rotas.
Gerenciador de Recursos	Gerenciar informações sobre os recursos da rede.
Interface de Plano de Controle	Comunicar com o plano de controle da rede.

Uma funcionalidade específica de serviços L1VPN é a gerência de informações de membros. Não é necessário apenas identificar quais membros pertencem a cada VPN. Também muito importantes, são mecanismos para identificar as propriedades desses membros (CEs de uma mesma VPN) e dos seus enlaces com a rede de transporte, assim como mecanismos para o compartilhamento automático dessas informações entre os elementos da rede de transporte. Essas informações são utilizadas quando do estabelecimento de conexões entre os CEs. Esses mecanismos permitem a descoberta automáticas de membros de uma VPN e de informações sobre esses membros (*VPN Membership Auto-Discovery* [44]).

Existem propostas em desenvolvimento que discutem implementar essas funções de forma distribuída, através de extensões a protocolos de roteamento tradicionais em redes IP [37, 7]. Tais mecanismos permitem que os elementos de borda da rede de transporte (PEs) compartilhem informações sobre os membros de uma VPN (CEs). Desta maneira, quando um novo CE é conectado à rede de transporte, sua identidade e propriedades são coletadas pelo respectivo PE e compartilhadas de forma automática com os outros elementos da rede do provedor. Assim, os outros membros da mesma VPN tomarão conhecimento do novo membro a partir das informações obtidas dos PEs aos quais eles estão conectados, e então, novas conexões poderão ser estabelecidas com o novo membro.

As políticas para gerência desses serviços podem ser elaboradas tanto pelos administradores da rede do provedor quanto pelos clientes do serviço. As políticas definidas pelo provedor podem ser gerais ou específicas para algum serviço. A fim de gerenciar o nível de controle oferecido aos clientes, o provedor deve avaliar e validar as políticas definidas pelos clientes. Mecanismos de prioridades também podem ser bastante úteis para garantir que as políticas definidas pelos clientes não sobreponham as definidas pelo provedor. A espe-

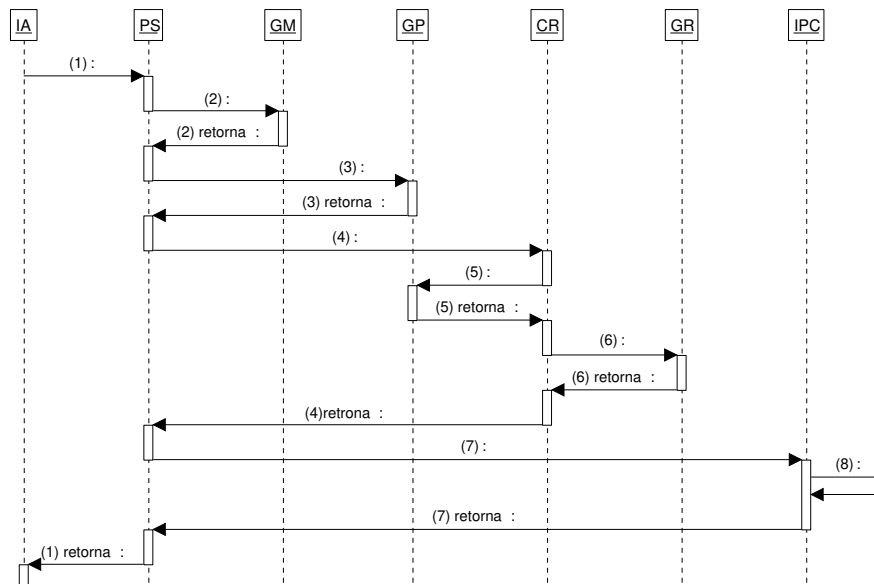


Figura 7.2: Interações entre os módulos para estabelecimento de conexão.

cificação do modelo de informações também tem papel importante nesta questão, como discutido no capítulo anterior. Esses aspectos devem ser cuidadosamente observados na implementação do gerenciador de políticas.

De fato, a gerência independente de cada serviço L1VPN resulta da definição de diferentes conjuntos de políticas que representem os objetivos e requisitos de cada cliente. O suporte à gerência baseada em políticas é efetivamente implementado ao se condicionar a operação dos módulos Provedor de Serviços, Controlador de Roteamento e Gerenciador de Recursos às decisões de políticas realizadas pelo Gerenciador de Políticas. Dessa maneira, as tarefas desempenhadas por aqueles módulos estão sujeitas às regras determinadas pelas políticas. Além disso, o gerenciador de políticas pode executar ações de reconfiguração definidas pelas políticas em função da ocorrência de determinados eventos. Este é o caso de alterar o funcionamento daqueles módulos, por exemplo, em virtude de alterações no estado da rede ou por causa de um novo requisito de um cliente.

Para esclarecer melhor o funcionamento dos módulos da arquitetura, é apresentado a seguir um exemplo de como os módulos interagem para o estabelecimento de uma conexão L1VPN. A Figura 7.2 apresenta um diagrama de seqüência que ilustra as interações entre os módulos. As etapas são descritas a seguir:

1. Um cliente envia uma requisição de conexão através da Interface de Acesso. Possivelmente, o cliente tem que passar por um procedimento de autenticação e verificação de permissão de acesso. A interface então encaminha a requisição ao módulo com-

petente para a tarefa: o Provedor de Serviços. O cliente será então notificado do resultado de sua requisição ao fim do procedimento.

2. O Provedor de serviços consulta o Gerenciador de Membros para verificar se os CEs de origem e destino da requisição de conexão pertencem à mesma VPN.
3. O Provedor de Serviços então requisita ao Gerenciador de Políticas decisões de gerência relacionadas com o serviço L1VPN em questão. Estas decisões servem para configurar os processos de admissão e estabelecimento da conexão.
4. Uma vez admitida a conexão, uma requisição é feita ao Controlador de Roteamento solicitando o cálculo de uma rota para a conexão. No fim desta tarefa, que ainda depende das duas etapas seguintes, uma rota é retornada ao Provedor de Serviços.
5. Os procedimentos para cálculo de rotas também estão sujeitos a decisões de políticas. Portanto o Controlador de Roteamento requisita ao Gerenciador de Políticas as decisões pertinentes.
6. O cálculo da rota depende também das informações de disponibilidade de recursos. Essas informações são então obtidas do Gerenciador de Recursos.
7. Por fim, o Provedor de Serviço solicita ao plano de controle o estabelecimento da conexão através da rede do provedor. Esta solicitação consiste em uma requisição enviada ao PE de ingresso conforme a rota calculada.
8. A conexão na rede do provedor é estabelecida através do mecanismo de sinalização do plano de controle, representado de forma simplificada no passo (8).

7.2 Cenários de Aplicação

Os cenários de aplicação da arquitetura proposta são descritos em termos de dois tipos de interface do serviço L1VPN: interface de gerência (caso do modelo baseado em gerência) e interface de controle (caso dos modelos baseados em sinalização e roteamento). Estes cenários caracterizam, respectivamente, as abordagens centralizada e distribuída. Em ambos os cenários, assume-se que a rede do provedor implementa um plano de controle GMPLS. Neste contexto, o estabelecimento dinâmico de conexão pode ser realizado por meio do mecanismo de sinalização GMPLS RSVP-TE [4] e a descoberta automática de topologia, através do mecanismo de roteamento GMPLS OSPF-TE [24].

As redes clientes e o provedor podem desempenhar diferentes tarefas em função da abordagem considerada. O nível de controle sobre o serviço L1VPN, atribuído a cada

cliente, depende de dois fatores. O primeiro corresponde à abordagem utilizada (centralizada ou distribuída). A abordagem centralizada limita as funções realizadas pelas redes clientes, em relação aos mecanismos de roteamento e recuperação de conexões. Isso porque esse modelo caracteriza um cenário *overlay*, no qual redes clientes e provedor apresentam planos de controle distintos, ou seja, não existe interação de sinalização ou roteamento entre cliente e provedor através de protocolos do plano de controle.

No caso da abordagem distribuída, pode haver interação entre as redes clientes e o provedor através do plano de controle. Desta forma, as redes clientes podem realizar funções de sinalização e roteamento. Por exemplo, o cálculo de rotas pode ser realizado pelas redes clientes, pois elas podem obter informações da topologia da rede do provedor através de trocas de mensagens de roteamento entre os planos de controle da rede cliente e da rede do provedor. Em um caso extremo, onde as redes clientes e a rede do provedor pertencem a uma mesma organização, pode existir um plano de controle comum entre essas redes (o que caracteriza um modelo *peer*). Um exemplo deste caso é a aplicação de serviços L1VPN na implementação de *multi-service backbone* (conforme explicado no Capítulo 3).

Um segundo fator que determina o nível de controle atribuído aos clientes é o conjunto de políticas que esses clientes podem especificar, conforme permitido pelo provedor. Por exemplo, mesmo no caso da abordagem centralizada, onde não há interação de controle de roteamento entre redes clientes e o provedor, pode ser permitido aos clientes especificar políticas para configuração do mecanismo de roteamento implementado pelo provedor.

7.2.1 Abordagem Centralizada

A Figura 7.3 ilustra o cenário de aplicação de serviço L1VPN considerando o modelo baseado em gerência. Neste caso, o provisionamento do serviço é baseado em uma arquitetura híbrida que combina funções distribuídas e centralizadas. As funções de controle são distribuídas, pois são realizadas pelo plano de controle GMPLS da rede do provedor. O mecanismo de sinalização distribuído e a descoberta automática de topologia aumentam a escalabilidade e reduzem o tempo de recuperação de conexões.

No entanto, as funções de gerência são centralizadas e podem ser implementadas em um sistema específico para gerência de serviços L1VPN ou integradas ao sistema de gerência de rede do provedor. Neste cenário, as funções de gerência centralizadas incluem as funções da arquitetura proposta. Assim, os módulos descritos são implementados no sistema de gerência de serviços L1VPN.

Para estabelecer uma conexão, o cliente envia uma requisição de conexão ao sistema de gerência. Depois de processar a requisição, o sistema comunica com o plano de controle, em nome do cliente, para estabelecer uma conexão através da rede do provedor, entre

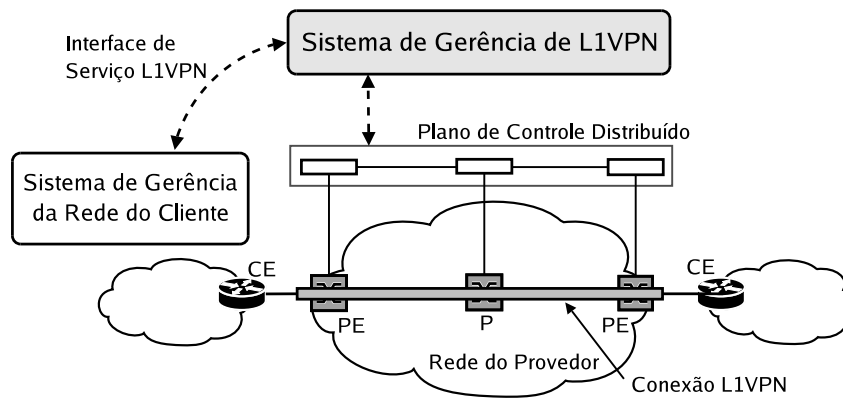


Figura 7.3: Serviço L1VPN baseado em interface de gerência.

os PEs correspondentes. Assim, calculada a rota para a conexão requisitada, o sistema de gerência solicita ao PE de ingresso o estabelecimento da conexão. Estabelecida a conexão, os CEs podem criar adjacências de roteamento sobre tal conexão, caracterizando um cenário *overlay*. A conexão requisitada pelo sistema de gerência e estabelecida pelo plano de controle representa então uma SPC (*Soft Permanent Connection*).

7.2.2 Abordagem Distribuída

Em um segundo cenário, com a interface do serviço no plano de controle, muitas das funcionalidades da arquitetura proposta podem ser implementadas de forma distribuída. Este modelo aumenta o nível de escalabilidade e robustez no provisionamento de serviços L1VPN. No entanto, ele representa uma solução de longo prazo, uma vez que extensões a protocolos e mesmo novas soluções ainda são necessárias.

As funções de gerência de informações de membros, controle e admissão de conexões, cálculo de rotas e gerência de recursos, antes implementadas no sistema centralizado, podem agora ser realizadas nos elementos de borda (PEs). Neste caso, as funções de gerência e compartilhamento automático de informações de membros podem ser realizadas por meio de um mecanismo de *VPN Membership Auto-Discovery*, por exemplo, utilizando-se o protocolo BGP, como descrito em [37], ou o protocolo OSPF [7].

No entanto, algumas funções, como gerência de políticas, permanecem centralizadas. Neste caso, existe o papel de um *Servidor de Políticas*, como ilustrado na Figura 7.4. Este servidor inclui um elemento de gerência de políticas e executa as funcionalidades de um PDP. Ele é logicamente centralizado e pode ser implementado em um sistema de gerência centralizado. Neste contexto, o controle e a gerência são governados pelas políticas pois as operações dos elementos PEs estão sujeitas às regras de políticas definidas no servidor.

Além disso, neste cenário de aplicação, os elementos PE precisam implementar as

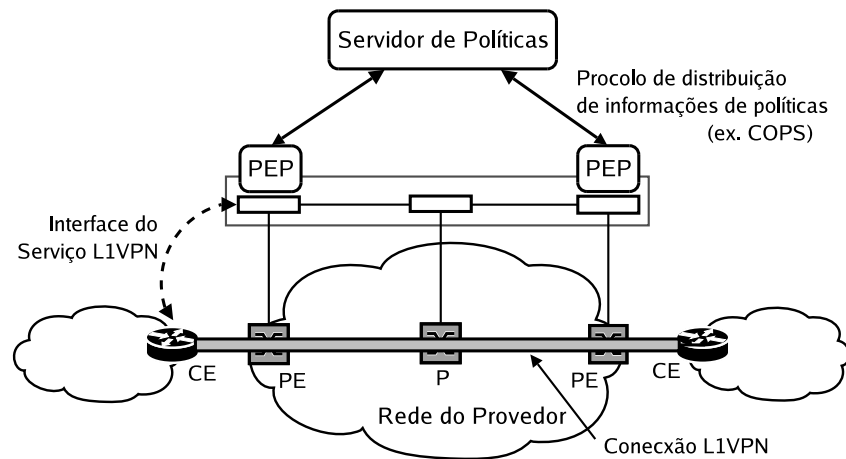


Figura 7.4: Serviço L1VPN baseado em interface de controle.

funcionalidades de um PEP a fim de suportar o mecanismo de gerência baseada em políticas. Os PEs se comunicam com o servidor de políticas para requisitar decisões de gerência. Para suportar este processo, deve ser utilizado um protocolo para a comunicação de informações de políticas, como o COPS (apresentado na Seção 4.3). Este protocolo também pode ser usado no modelo *provisioning* para prover o PE com as informações de configuração dos serviços L1VPN.

Mesmo no caso da abordagem distribuída, o provedor pode ainda manter um sistema de gerência centralizado, a fim de prover os clientes com funções de monitoramento, como obtenção de informações sobre desempenho e falhas do serviço L1VPN correspondente.

Parte III

Avaliação e Conclusão

Capítulo 8

Implementação e Avaliação

O objetivo da avaliação da arquitetura proposta não é realizar uma análise quantitativa dos resultados numéricos, mas sim, discutir os diferentes efeitos das políticas sob diferentes perspectivas. Na primeira seção, é apresentada a implementação de um protótipo da arquitetura proposta. Na Seção 8.2, é apresentado um estudo de caso no qual são avaliados dois cenários simulados.

8.1 Implementação

Foi implementado um protótipo para validar a arquitetura proposta, como ilustrado na Figura 8.1. A implementação considera o modelo de serviço baseado em gerência, como descrito em um dos cenários de aplicação da arquitetura, no capítulo anterior. Assim, funções como cálculo de rotas e controle de conexões são centralizadas. Entretanto, o protótipo foi implementado como um sistema distribuído utilizando-se duas tecnologias: *Java Remote Method Invocation* (Java RMI) e *Web Services*.

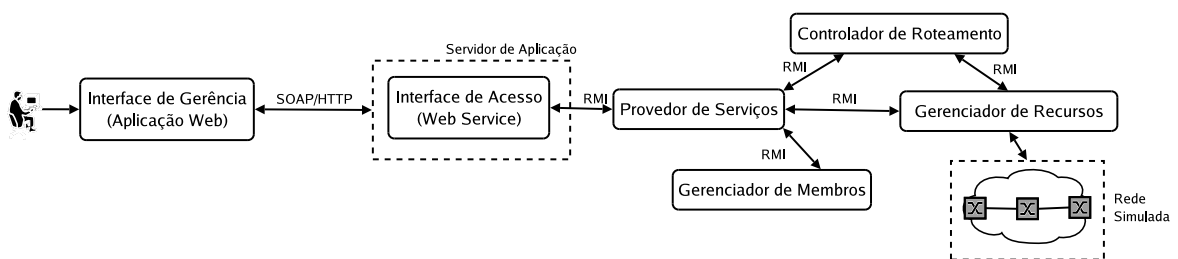


Figura 8.1: Estrutura do protótipo implementado.

Os módulos principais da arquitetura foram implementados como objetos remotos utilizando-se a tecnologia Java RMI, que oferece uma plataforma robusta para a implementação de aplicações distribuídas. O módulo Provedor de Serviços possui dois sub-

módulos responsáveis pelo controle de admissão e pelo controle de conexões. O módulo Controlador de Roteamento implementa o algoritmo de menor caminho de Dijkstra para o cálculo de rotas. O módulo Monitor de Serviços não foi implementado.

Nesta implementação, as políticas são definidas de forma estática (em tempo de compilação). Isso significa que as regras das políticas são implementadas nos módulos correspondentes. Não foi implementado um módulo Gerenciador de Políticas independente que permitisse o controle dinâmico de políticas (em tempo de execução). Por exemplo, políticas de controle de admissão são implementadas diretamente no código do Provedor de Serviços.

Por questões de flexibilidade e interoperabilidade, o módulo Interface de Acesso foi implementado como um *Web Service*. A tecnologia *Web Services* suporta o desenvolvimento de sistemas distribuídos facilitando a interoperabilidade na comunicação entre aplicações através do uso de protocolos da Internet e de especificações baseadas em XML (*Extensible Markup Language*). Foi utilizado o servidor de aplicações Apache Tomcat e como implementação do protocolo SOAP (*Simple Object Access Protocol*), foi utilizado o Axis, também da Fundação Apache. O SOAP é um protocolo baseado em XML para a troca de informações em um ambiente distribuído. Ele é utilizado para enviar requisições e receber respostas de *Web Services*.

A interface do protótipo (Interface de Gerência, na figura) foi implementada como uma aplicação *Web*. Através dela o usuário pode acessar as funções do sistema. Para atender as requisições, essa interface invoca o método correspondente do *Web Service* que implementa o módulo Interface de Acesso. A Figura 8.2 ilustra uma tela da interface do protótipo. Como pode ser observado no menu do lado esquerdo da figura, as funções disponíveis são:

- Instanciar serviços (*L1VPN Service*);
- Gerenciar informações sobre os membros de uma VPN (*Membership Control*);
- Editar políticas¹ (*Policies*);
- Estabelecer ou remover conexões (*Connection Control*).

A tela apresentada nessa figura mostra também a interface para controle de informações de membros. As funções disponíveis permitem adicionar ou remover um membro de um serviço L1VPN, ou ainda listar todos os membros de um serviço.

Na Figura 8.3 é ilustrada a tela com funções para controle de conexões. A interface inclui funções para o usuário requisitar o estabelecimento ou a remoção de conexões. Para

¹A aplicação permite editar ações para configuração de modelo de alocação de recursos e de métrica para peso de enlaces

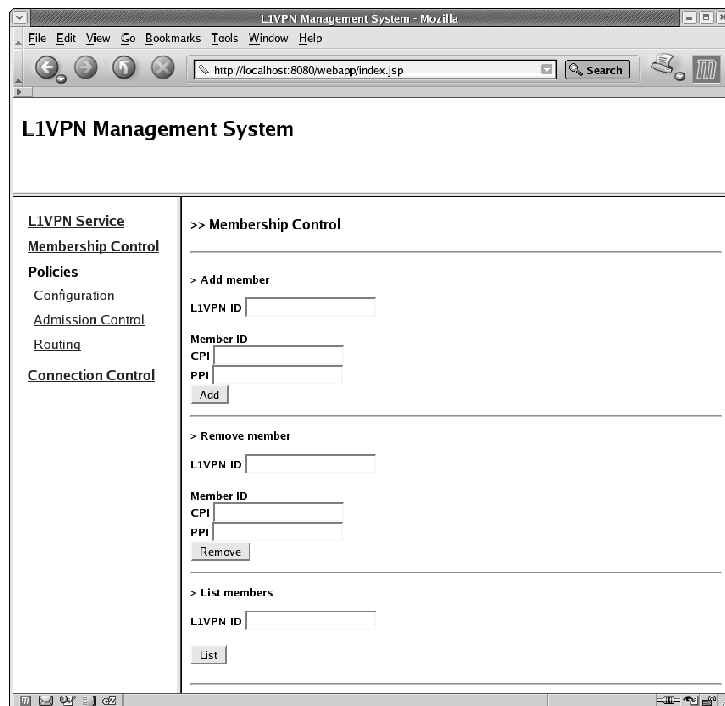


Figura 8.2: Tela de gerência de membros.

estabelecer uma conexão é preciso fornecer o identificador do serviço VPN para o qual será criada a conexão, além dos identificadores dos nós de ingresso e egresso.

A tela apresentada na Figura 8.4 demonstra a edição de uma política. Neste caso, a condição define para qual L1VPN serão aplicadas as ações definidas. Estas ações consistem em determinar o modelo de alocação de recursos e o critério para definição do peso de enlaces para o algoritmo de cálculo de rotas. No caso de escolha do modelo dedicado, ainda é possível determinar a quantidade de recursos que serão reservados para o serviço, em termos do número de comprimentos de onda em cada enlace da rede. Observe que a interface de gerência permite somente a edição dos parâmetros das políticas de configuração, mas não a edição de novas políticas. Novas políticas precisam ser implementadas diretamente no código do protótipo.

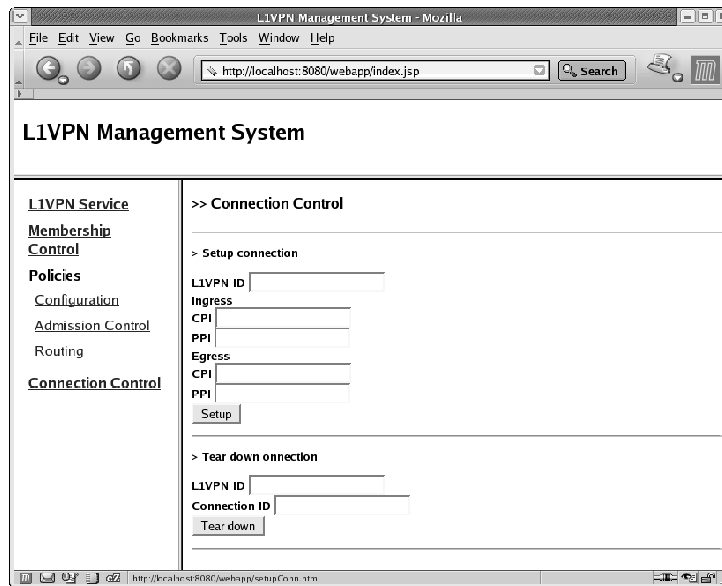


Figura 8.3: Tela de controle de conexão.

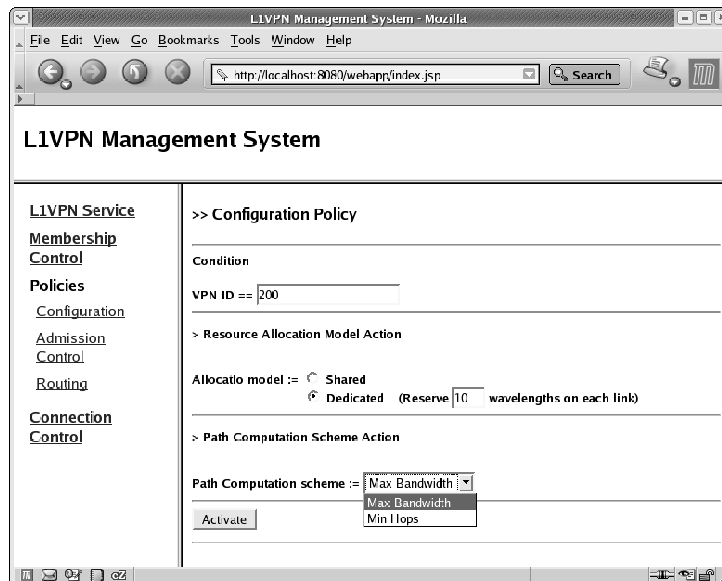


Figura 8.4: Tela de edição de política.

O módulo Gerenciador de Recursos implementa uma infra-estrutura de rede transporte óptica, cuja topologia é apresentada na Figura 8.5. Esta topologia representa a rede NSFNet com 14 nós. O Gerenciador de Recursos também mantém informações sobre a disponibilidade de recursos, por exemplo, quais comprimentos de onda estão disponíveis em cada enlace.

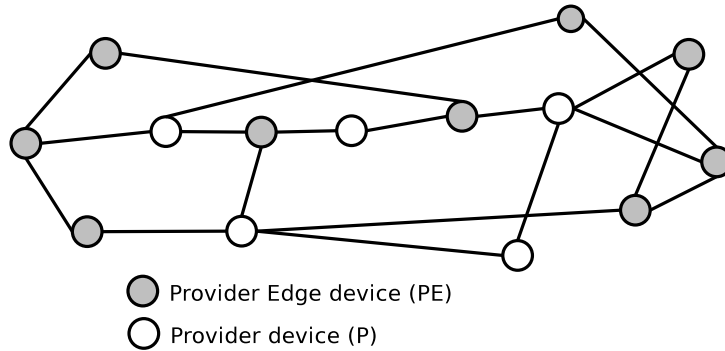


Figura 8.5: Topologia de rede simulada.

O estabelecimento de uma conexão na rede óptica consiste em determinar uma rota e alocar um comprimento de onda disponível em cada enlace do caminho do nó origem até o destino. Um mecanismo de sinalização foi implementado através de um agente de controle que é incorporado a cada nó. Esse mecanismo de controle inclui funções para alocar ou liberar um comprimento de onda em um enlace. Esses agentes são responsáveis por alocar recursos durante o estabelecimento de uma conexão, informando o gerenciador de recursos sobre a alocação. Uma mensagem de requisição é enviada ao nó de origem juntamente com informações sobre a rota da conexão. Cada nó então comunica com o próximo nó na rota para reservar um comprimento de onda entre eles, até o nó destino.

8.2 Estudo de Caso

Foi realizado um estudo de caso para avaliar os diferentes efeitos de políticas. Foram consideradas políticas de configuração do modelo de alocação de recursos e do mecanismo de cálculo de rotas. Sob a perspectiva do cliente foi considerada a taxa de bloqueio de conexões. Uma requisição de conexão é bloqueada quando não há recursos (comprimentos de onda) disponíveis para o estabelecimento da conexão. Por outro lado, do ponto de vista do provedor, foi analisada a taxa de utilização dos recursos da rede, em termos do número total de comprimentos de onda alocados para as conexões.

8.2.1 Ambiente de Simulação

O estudo de caso foi realizado através da implementação de um ambiente de simulação, cuja estrutura é apresentada na Figura 8.6. Este ambiente de simulação foi implementado a partir dos módulos desenvolvidos para o protótipo. No entanto, diferente do protótipo, ele não foi implementado como um sistema distribuído. Assim, os módulos do ambiente de simulação são executados no mesmo computador, como partes de um único sistema que simula a dinâmica do estabelecimento de conexões de serviços L1VPN em uma rede óptica. A topologia e o mecanismo de sinalização são os mesmos utilizados no protótipo. Esse ambiente permite simular cenários com diversos serviços L1VPN nos quais os clientes requisitam conexões de forma concorrente. Isso foi implementado por meio de programação concorrente, de forma que cada cliente é representado por uma *thread* no sistema.

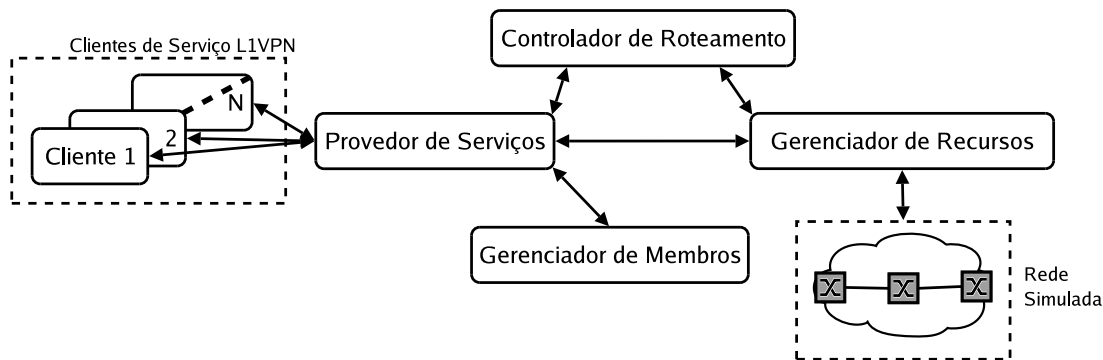


Figura 8.6: Estrutura do ambiente de simulação.

Em cada execução de uma simulação, cada cliente solicita um total de 2500 conexões. Como o controle de conexões é centralizado, evita-se o problema de conflito na alocação de recursos. As requisições de conexão são atendidas por ordem de chegada, de forma síncrona, no módulo Provedor de Serviços. Foi considerado que para cada cliente de um serviço L1VPN existe um CE conectado a cada PE (conforme a topologia da rede do provedor apresentada na Figura 8.5). Portanto, cada serviço L1VPN possui 9 membros. Além disso, as conexões são estabelecidas entre PEs. Dessa maneira, a saturação de recursos nos enlaces entre CE e PE não é um problema e apenas os recursos do provedor são contabilizados nos resultados.

Os pares origem e destino para as conexões são escolhidos segundo uma distribuição aleatória uniforme entre o conjunto de CEs membros de uma VPN. A taxa de requisição de conexões segue uma distribuição de Poisson², como é usual nos trabalhos de simulação em redes ópticas com roteamento por comprimento de onda (*wavelength-routed optical networks*). Neste caso, foi considerada uma média igual a 100. A distribuição de Poisson

²O processo de Poisson é implementado conforme apresentado em [20, pag. 496]

é apropriada para modelar o número de chegadas em dado intervalo de tempo, principalmente considerando diversas fontes independentes [20, pag. 496].

O tempo de duração de uma conexão (até que os recursos sejam liberados) segue uma distribuição exponencial³, com média igual a 1200. O intervalo entre requisições segue uma distribuição exponencial cuja média é dada em função da taxa de requisições. A distribuição exponencial é apropriada para modelar intervalos de tempo entre eventos sucessivos e intervalos de tempo de atendimento de um serviço (*service time*) [20, pag. 489]. Os resultados numéricos são a média de 100 iterações das simulações.

8.2.2 Primeiro Cenário

Neste cenário, o objetivo é avaliar o efeito das políticas sob a perspectiva do cliente. Neste caso, é avaliado o efeito das políticas na taxa de bloqueio de conexões. A simulação envolve 4 serviços (L1VPN 0-3) e cada enlace da rede possui 32 comprimentos de onda. Para esta avaliação são consideradas duas classes para os serviços L1VPN: serviços de alta prioridade e de baixa prioridade. As duas políticas a seguir são utilizadas para gerenciar um serviço de acordo com sua classe:

1. Se a VPN é de alta prioridade, então:
 - A alocação de recursos segue o modelo dedicado. Um subconjunto dos recursos do provedor é dedicado para a VPN;
 - No cálculo de rotas, devem ser considerados apenas os recursos dedicados para a VPN;
 - O critério para determinar a rota é escolher os enlaces com maior número de comprimentos de onda disponíveis.
2. Se a VPN é de baixa prioridade, então:
 - A alocação de recursos segue o modelo compartilhado. A VPN disputa com outras VPNs a alocação de recursos compartilhados;
 - No cálculo de rotas, apenas os recursos compartilhados são considerados;
 - O critério para determinar a rota é escolher o menor caminho (em número de *hops*).

Neste cenário, para um serviço considerado como de alta prioridade, são reservados 10 comprimentos de onda em cada enlace. Para implementar o cálculo de rotas priorizando

³A distribuição exponencial é implementada conforme apresentado em [20, pag. 489]

os enlaces com mais comprimentos de onda disponíveis, assumimos que o custo do enlace é definido por $\frac{1}{w}$, onde w é o número de comprimentos de onda disponíveis (não alocados) no enlace. Por outro lado, quando o critério é o menor caminho, assumimos que o peso dos enlaces é igual a 1.

Primeiro a simulação foi realizada com todos os serviços definidos como de baixa prioridade. A Figura 8.7(a) apresenta a taxa de bloqueio da L1VPN 0 e a taxa de bloqueio média das outras L1VPNs. A Figura 8.7(b) mostra a alteração nesses valores quando a simulação foi repetida, agora com a L1VPN 0 definida como serviço de alta prioridade.

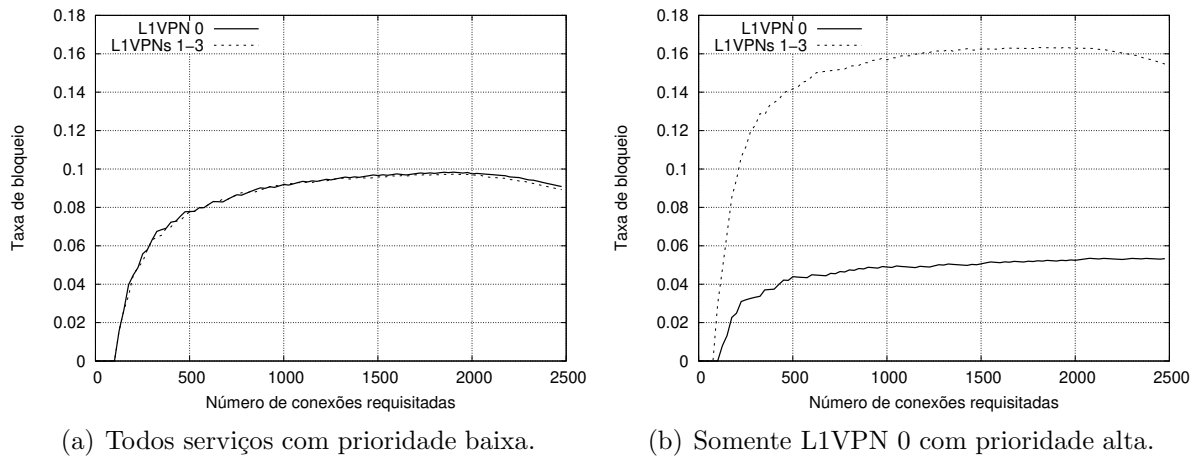


Figura 8.7: Taxa de bloqueio de conexões.

Os resultados demonstram uma queda na taxa de bloqueio para o serviço L1VPN 0, como efeito das ações de gerência definidas pela segunda política. De fato, a política que configura serviços de alta prioridade foi definida de forma a privilegiar tais serviços com recursos dedicados suficientes para diminuir a taxa de bloqueio. Isso demonstra como as políticas podem ser utilizadas para configurar diferentes classes de serviços.

Um caso mais interessante é apresentado na Figura 8.8. Ela demonstra como políticas podem ser usadas para definir como o sistema pode reagir a alterações nas condições da rede. O resultado mostra uma diminuição na taxa de bloqueio da VPN 0 quando é ativada uma política que altera a prioridade dessa VPN. Neste caso a política é ativada quando o número de conexões requisitadas atinge um certo valor. No entanto, condições mais elaboradas podem ser definidas. Por exemplo, para responder à ocorrência de eventos específicos ou para alterar a configuração de serviços de alta prioridade, quando é ultrapassado um limiar de nível de degradação de um parâmetro de qualidade de serviço.

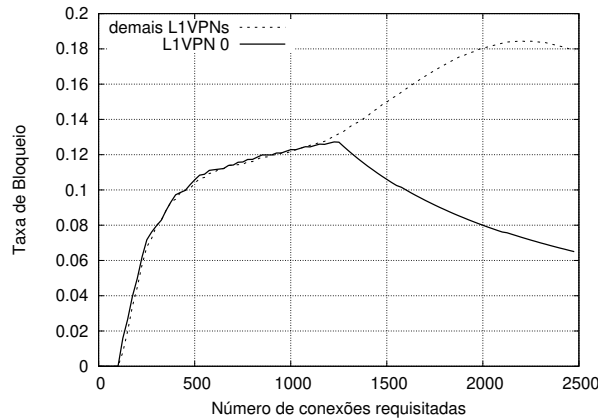


Figura 8.8: Efeito de ativação da política.

8.2.3 Segundo Cenário

Neste segundo cenário, o objetivo é avaliar o efeito das políticas de configuração do modelo de alocação de recursos sobre a taxa de utilização de recursos da rede de transporte. Este cenário considera 8 serviços L1VPN e 64 comprimentos de onda em cada enlace. As simulações forem realizadas com diferentes taxas de requisição de conexões, primeiro com todas as VPNs configuradas para utilizar o modelo compartilhado e depois o modelo dedicado. No caso do modelo dedicado, os recursos são divididos igualmente entre as L1VPNs.

Para baixas taxas de requisição, não há diferenças significativas entre os modelos de alocação (Figura 8.9(a)). Por outro lado, para taxas mais altas, o modelo compartilhado garante melhor utilização dos recursos (Figura 8.9(b)). A variação na taxa de requisição é obtida variando-se a média do processo de Poisson que define a taxa. Assim, uma taxa de requisição igual a 0,5 consiste em adotar o valor da média como metade do valor da média base (igual a 100).

Novamente, a política de configuração que define o modelo como dedicado implicou num resultado geral desfavorável para o provedor, considerando a média de utilização dos recursos da rede. No entanto, essa política foi utilizada em outros casos para melhorar a taxa média de bloqueio para serviços específicos.

8.2.4 Discussão Final

O efeito das políticas pode ser diferente quando consideradas em conjunto ou sob diferentes prioridades. Por exemplo, as políticas consideradas no primeiro cenário definiram o modelo de alocação de recursos como dedicado, a fim de melhorar o desempenho de um serviço específico (em termos da taxa de bloqueio). No entanto, o modelo dedicado pode

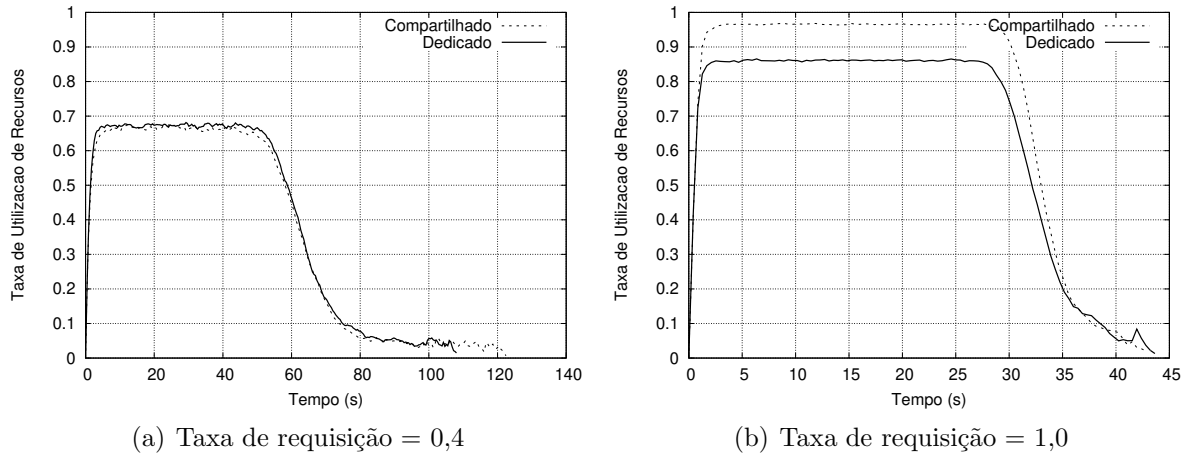


Figura 8.9: Taxa de utilização de recursos da rede do provedor.

levar a uma degradação da performance geral da rede ou dos serviços. A Figura 8.10 compara o desempenho dos modelos de alocação em relação à taxa média de bloqueio das VPNs, para diferentes taxas de requisição de conexões. Neste caso, foram considerados 8 serviços L1VPN. A variação da taxa de requisição é análoga ao modo como foi feita no segundo cenário, utiliza-se uma fração da média base na distribuição de Poisson.

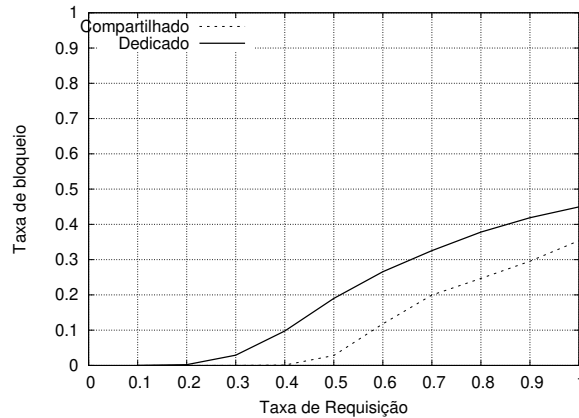


Figura 8.10: Taxa de bloqueio: modelo compartilhado x dedicado.

Neste caso, primeiro foram realizadas simulações com todos os serviços configurados para utilizar o modelo dedicado sob diferentes taxas de requisição. Depois, as simulações foram repetidas com os serviços configurados para utilizar o modelo compartilhado. No caso do modelo dedicado, os recursos são divididos igualmente entre as L1VPNs. Neste avaliação, o modelo dedicado foi menos eficiente, considerando a taxa de bloqueio média dos serviços. No entanto, nos casos anteriores, o modelo dedicado foi utilizado para

melhorar a taxa de bloqueio para serviços específicos.

Entretanto, o modelo dedicado pode apresentar resultados melhores que o compartilhado, dependendo de outros fatores como o modelo de serviço L1VPN. No modelo dedicado, os clientes podem receber informações detalhadas sobre a topologia e o estado dos recursos para eles alocados. Neste caso, um algoritmo de cálculo de rota do cliente pode otimizar a utilização dos recursos. O trabalho apresentado em [48] propõe algoritmos de alocação de recursos e descreve um cenário de redes multi-camadas (*multilayer*) no qual o modelo dedicado supera o compartilhado em termos da taxa de bloqueio de conexões. A justificativa é que a eficiência de algoritmos otimizados supera as desvantagens do fato de os recursos não serem compartilhados.

O segundo cenário também demonstra os benefícios e desvantagens de uma mesma política quando considerada sob diferentes perspectivas. As simulações demonstram que uma configuração de modelo dedicado resulta em desempenho inferior na taxa de utilização dos recursos da rede. No entanto, a política pode ser aplicada para melhorar o desempenho de um conjunto específico de serviços.

Os cenários simulados são importantes para demonstrar como o provedor pode definir políticas para configurar diferentes classes de serviços L1VPN. E também ilustra como políticas podem ser utilizadas para reconfigurar os serviços em resposta a mudanças nas condições da rede ou em resposta a níveis insatisfatórios de determinado parâmetro de desempenho do serviço (por exemplo, taxa de bloqueio de conexões). O provedor precisa observar os diferentes efeitos das políticas para controlar o nível de gerência atribuído aos clientes e para definir o modelo de serviços L1VPN oferecido e sua estratégia de gerência para esses serviços.

Capítulo 9

Conclusão

O controle dinâmico de conexões, proporcionado por um plano de controle distribuído, possibilita o provisionamento de serviços VPN em redes de camada 1. Desta forma, esse serviço permite que vários clientes compartilhem uma rede de transporte comum. Os clientes podem alocar recursos da rede do provedor para interconectar suas redes e utilizar esses recursos como uma infra-estrutura de transporte dedicada, sem a necessidade de implantar e operar sua própria rede. Existe uma expectativa de que L1VPN será um dos mais importantes serviços nas redes de próxima geração.

Muitos trabalhos têm discutido como prover serviços L1VPN considerando questões relacionadas com o plano de controle, como sinalização e endereçamento. Porém faltam propostas sobre como atender os requisitos de gerência desses serviços. Neste trabalho, foi proposta uma arquitetura para a gerência de serviços L1VPN. A metodologia foi utilizar a abordagem de Gerência Baseada em Políticas.

A arquitetura proposta não define mecanismos de sinalização ou roteamento. Considera-se uma rede de transporte com um plano de controle, como a arquitetura GMPLS, que implementa essas funcionalidades. Com a arquitetura proposta, o provedor pode oferecer aos clientes um certo nível de controle e gerência sobre seus serviços L1VPN de forma independente. Isso é obtido por permitir a cada cliente definir, de forma supervisionada, um conjunto de políticas que gerenciam a configuração e operação do seu serviço.

No contexto da abordagem PBM, foram definidas classes de políticas para gerência de serviços L1VPN e também foi discutido como o *framework* de políticas da IETF pode ser utilizado na gerência desses serviços. Além disso, foi apresentada uma alternativa para a representação de políticas para serviços L1VPN, por meio da definição de um modelo de informação e do uso de tecnologias XML.

A definição das classes de políticas foi importante para organizar e ilustrar diversos tipos de políticas que podem ser utilizadas para serviços L1VPN. Esta etapa foi importante para esclarecer a importância de um Modelo de Informações de Políticas padronizado,

principalmente por questões de interoperabilidade e como forma de controlar as políticas elaboradas pelos clientes. Uma abordagem interessante é definir um modelo específico para os clientes.

A representação de políticas também desempenha papel fundamental e, portanto, é necessária uma decisão cuidadosa sobre a linguagem e o modelo de dados utilizados para representar as políticas. Apesar de linguagens específicas para representação de políticas, como PONDER, existe uma tendência em se utilizar tecnologias XML para esse fim. Uma contribuição nesse sentido foi a proposta de modelo de informações, baseado nas classes propostas, e do uso de XML para representar as políticas.

Foi descrito o modelo funcional e cenários de uso da arquitetura proposta. Vale ressaltar a importância para um provedor de se analisar os benefícios e custos das abordagens híbrida e distribuída no uso da arquitetura. A primeira opção apresentada é menos eficiente em termos de escalabilidade e desempenho, devido a problemas inerentes de alguns mecanismos centralizados. Porém é uma solução atraente a curto prazo, pois dispensa a tarefa complexa de estender ou implementar alguns mecanismos do plano de controle.

A descrição dos cenários de aplicação foi importante para caracterizar as diferentes abordagens para implementação das funcionalidades do serviço L1VPN, assim como os fatores que determinam o nível de controle atribuído aos clientes do serviço. Esses fatores são o modelo da arquitetura utilizado (centralizado ou distribuído) e o subconjunto de políticas que podem ser especificadas pelos clientes, conforme definido pelo provedor.

A implementação de um protótipo demonstrou a viabilidade da arquitetura e as vantagens de se combinar diferentes tecnologias para implementação de sistemas distribuídos. Neste caso, foi utilizado Java RMI, para um maior grau de robustez e desempenho aos módulos internos da arquitetura, e a tecnologia *Web Services*, para a interface de gerência, o que contribuiu com um mecanismo flexível para acesso ao sistema, baseado em protocolos amplamente utilizados na Internet.

Além disso, um estudo de caso foi realizado para avaliar a arquitetura e as implicações do uso de políticas. A partir de cenários simulados foram avaliadas políticas de configuração de modelo de alocação de recursos e mecanismos de cálculo de rotas. Este estudo foi importante para demonstrar como a arquitetura baseada em políticas pode ser utilizada para configurar diferentes classes de serviços L1VPN e para reconfigurar os serviços em função de degradação de desempenho ou alterações no estado da rede.

Os cenários apresentados também mostram como as políticas podem ser aplicadas e seus diferentes efeitos, quando se considera diferentes perspectivas, de provedor e cliente. As políticas que melhoram o desempenho de um serviço podem afetar o desempenho de outros serviços. Assim como, uma política que melhora o desempenho para um cliente, pode trazer desvantagens para o provedor, por exemplo, afetar a taxa de utilização dos recursos da rede.

De fato, as políticas definidas pelos clientes podem impactar fortemente no desempenho geral da rede do provedor. O provedor deve se preocupar com quais políticas um cliente pode especificar, pois uma política que traz benefícios para um cliente pode afetar o serviço de outro. Ou mesmo o efeito de uma política pode ser diferente sob diferentes condições da rede.

O alto nível de abstração das políticas e a própria complexidade da rede como um sistema dificultam a precisão em determinar os efeitos resultantes. Um provedor poderia utilizar-se de ferramentas de simulação e modelagem analítica para estimar efeitos de políticas, além de ter que controlar rigidamente as regras de políticas definidas pelos clientes. A resolução de conflitos entre políticas também é um problema comum que merece atenção. Neste caso, mecanismos de prioridades representam uma solução razoável para a maioria dos casos. Estes mecanismos também são importantes para limitar o nível de controle oferecido aos clientes.

9.1 Trabalhos Futuros

Este trabalho envolveu estudar diversos assuntos, alguns deles, como o próprio conceito de serviços L1VPN, ainda muito recentes ou não muito explorados na literatura. Em particular, foi abordado apenas o caso de serviços L1VPN em um único domínio administrativo.

Um importante trabalho futuro seria investigar o provisionamento de serviços L1VPN em múltiplos domínios. O compartilhamento de informações entre domínios envolve muitos desafios. Um trabalho neste sentido envolveria pesquisar como isso afeta funções como descoberta de informações de membro, endereçamento, roteamento das conexões VPN por domínios com diferentes planos de controle, entre outros. Outro aspecto fundamental é como prover Qualidade de Serviço neste contexto. Em geral, serviços VPN são utilizados para aplicações de missão crítica que exigem restrições rígidas em relação a requisitos de qualidade de serviço. O estabelecimento de conexões VPN sobre múltiplos domínios implicam em muitos desafios para atender a esses requisitos.

Além de serviços VPN sobre múltiplos domínios, outros importantes trabalhos futuros incluem:

- Estudar algoritmos e mecanismos baseados em políticas para alocação de recursos. O modelo de alocação é um fator crítico no estabelecimento de conexões. É importante desenvolver algoritmos e mecanismos de alocação de recursos no sentido de melhorar o desempenho do serviço em diferentes aspectos como taxa de bloqueio de conexões e taxa de utilização de recursos.

- Realizar uma avaliação abrangente das políticas para gerência de serviços L1VPN, considerando as diversas classes propostas.
- Implementar um protótipo da arquitetura de gerência proposta considerando uma abordagem distribuída e avaliar aspectos de desempenho e escalabilidade.

Referências Bibliográficas

- [1] D. Agrawal, Kang-Won Lee, and J. Lobo. Policy-based management of networked computing systems. *Communications Magazine, IEEE*, 43(10):69–75, 2005.
- [2] S. Baek, M. Jeong, J. Park, and T. Chung. Policy-based hybrid management architecture for IP-based VPN. In *Network Operations and Management Symposium, NOMS 2000*, pages 987–988, Abril 2000.
- [3] A. Banerjee, J. Drake, J.P. Lang, B. Turner, K. Kompella, and Y. Rekhter. Generalized Multiprotocol Label Switching: An Overview of Routing and Management Enhancements. *Communications Magazine, IEEE*, 39(1):144–150, 2001.
- [4] L. Berger. Generalized Multi-Protocol Label Switching (GMPLS) Signaling - Resource ReserVation Protocol-Traffic Engineering RSVP-TE) Extensions. IETF RFC 3473, Janeiro 2003.
- [5] L. Berger. Generalized Multi-Protocol Label Switching (GMPLS) Signaling Functional Description. IETF RFC 3471, Janeiro 2003.
- [6] G. Bernstein, B. Rajagopalan, and D. Saha. *Optical Network Control: Architecture, Protocols, and Standards*. Addison-Wesley, 2003.
- [7] I. Bryskin and L. Berger. OSPF based L1VPN auto-discovery. IETF Internet-Draft, “work in progress”, Fevereiro 2006.
- [8] C. Carvalho, E. Madeira, F. Verdi, and M. Magalhães. Policy-Based Fault Management for Integrating IP over Optical Networks. In *5th International Workshop on IP Operations and Management, IPOM 2005*, volume 3751 of *LNCS*, pages 88–97, Barcelona, Espanha, 2005.
- [9] M. Chamoun, R. Kilany, and A. Serhrouchni. A semantic active policy-based management architecture. In *IEEE Workshop on IP Operations and Management, IPOM 2004*, pages 224–232, Outubro 2004.

- [10] W. Changkun. Policy-based Network Management. In *International Conference on Communication Technology Proceedings, ICCT 2000*, volume 1, pages 101–105, 2000.
- [11] B. Daheb, W. Fawaz, O. Audouin, B. Berde, K. Chen, and G. Pujolle. Policy-based hybrid hierarchical optical networks. In *First International Conference on Broadband Networks Proceedings, BroadNets 2004*, pages 325–327, Outubro 2004.
- [12] N. Damianou, N. Dulay, E. Lupu, and M. Sloman. The Ponder Policy Specification Language. In *Policies for Distributed Systems and Networks: International Workshop Proceedings, POLICY 2001*, volume 1995 of *LNCS*, pages 18–38, Bristol, UK, 2001.
- [13] N. Damianou, N. Dulay, E. Lupu, M. Sloman, and T. Tonouchi. Tools for Domain-Based Policy Management of Distributed Systems. In *Network Operations and Management Symposium, 2002, NOMS 2002*, pages 203–217, 2002.
- [14] D. Durham. The COPS (Common Open Policy Service) Protocol. IETF RFC 2748, Janeiro 2000.
- [15] D. Fedyk and Y. Rekhter. Layer 1 VPN Basic Mode. IETF Internet-Draft, “work in progress”, Maio 2006.
- [16] W. Golab and R. Boutaba. Policy-driven automated reconfiguration for performance management in WDM optical networks. *Communications Magazine, IEEE*, 42(1):44–51, 2004.
- [17] X. Guo, K. Yang, A. Galis, X. Cheng, B. Yang, and D. Liu. A policy-based network management system for IP VPN. In *International Conference on Communication Technology Proceedings, ICCT 2003*, volume 2, pages 1630–1633, Abril 2003.
- [18] ITU. Layer 1 Virtual Private Network generic requirements and architecture elements. ITU-T Recommendation Y.1312, Setembro 2003.
- [19] ITU. Layer 1 Virtual Private Network service and network architectures. ITU-T Recommendation Y.1313, Julho 2004.
- [20] R. Jain. *The Art of Computer Systems Performance Analysis: Techniques for Experimental Design, Measurement, Simulation, and Modeling*. John Wiley & Sons, 1991.
- [21] K. Kompella and Y. Rekhter. IS-IS Extensions in Support of Generalized Multi-Protocol Label Switching (GMPLS). IETF RFC 4205, Outubro 2005.

- [22] K. Kompella and Y. Rekhter. Label Switched Paths Hierarchy with GMPLS Traffic Engineering. IETF RFC 4206, Outubro 2005.
- [23] K. Kompella and Y. Rekhter. OSPF Extensions in Support of Generalized Multi-Protocol Label Switching (GMPLS). IETF RFC 4203, Outubro 2005.
- [24] K. Kompella and Y. Rekhter. Routing Extensions in Support of Generalized Multi-Protocol Label Switching (GMPLS). IETF RFC 4202, Outubro 2005.
- [25] J. F. Kurose and K. W. Ross. *Computer Networking: A Top-Down Approach Featuring the Internet*. Addison Wesley, 3 edition, 2004.
- [26] J. Lang. Link Management Protocol (LMP). IETF RFC 4204, Outubro 2005.
- [27] K. L. E. Law and A. Saxena. Scalable Design of a Policy-Based Management System and its Performance. *Communications Magazine, IEEE*, 41(6):72–79, 2003.
- [28] M. Li. Policy-based IPsec management. *Network, IEEE*, 17(6):36–43, 2003.
- [29] N. Malheiros, E. Madeira, F. Verdi, and M. Magalhães. A Management Architecture for Layer 1 VPN Services. In *International Conference on Broadband Communications, Networks and Systems, BroadNets 2006*, San Jose, Estados Unidos, Outubro 2006.
- [30] N. Malheiros, E. Madeira, F. Verdi, and M. Magalhães. Uma Arquitetura Baseada em Políticas para Gerência de VPNs de Camada 1. In *Anais do XXIV Simpósio Brasileiro de Redes de Computadores, SBRC 2006*, Curitiba, Brasil, Maio/Junho 2006.
- [31] E. Mannie. Generalized Multi-Protocol Label Switching Architecture. IETF RFC 3945, Outubro 2004.
- [32] E. Mannie and D. Papadimitriou. Recovery (Protection and Restoration) Terminology for GMPLS. IETF Internet-Draft, “work in progress”, Abril 2005.
- [33] J. V. Millor and J. S. Fernandez. A network management approach enabling active and programmable Internets. *Network, IEEE*, 19(1):18–24, 2005.
- [34] B. Moore. Policy Core Information Model (PCIM) Extensions. IETF RFC 3460, Janeiro 2003.
- [35] B. Moore, E. Ellessen, J. Strassner, and A. Westerinen. Policy Core Information Model – Version 1 Specification. IETF RFC 3060, Fevereiro 2001.

- [36] H. Ould-Brahim. GVPN Services: Generalized VPN Services using BGP and GMPLS Toolkit. IETF Internet-Draft, “work in progress”, Agosto 2005.
- [37] H. Ould-Brahim, D. Fedyk, and Y. Rekhter. BGP-based auto-discovery for L1VPNs. IETF Internet-Draft, “work in progress”, Março 2006.
- [38] A. Pras, J. Schornwalder, and O. Festor. Guest editorial: XML-based management of networks and services. *Communications Magazine, IEEE*, 42(7):56–57, 2004.
- [39] E. Rosen, A. Viswanathan, and R. Callon. Multiprotocol Label Switching Architecture. IETF RFC 3031, Janeiro 2001.
- [40] J. Strassner. *Policy-Based Network Management: Solutions for the Next Generation*. Morgan Kaufmann, 2003.
- [41] G. Swallow, J. Drake, H. Ishimatsu, and Y. Rekhter. GMPLS User-Network Interface. IETF RFC 4208, Outubro 2005.
- [42] T. Takeda, R. Aubin, M. Carugi, I. Inoue, and H. Ould-Brahim. Framework and Requirements for Layer 1 Virtual Private Networks. IETF Internet-Draft, “work in progress”, Março 2006.
- [43] T. Takeda, D. Brungard, A. Farrel, H. Ould-Brahim, and D. Papadimitriou. Applicability analysis of GMPLS protocols to Layer 1 Virtual Private Networks. IETF Internet-Draft, “work in progress”, Março 2006.
- [44] T. Takeda, D. Brungard, D. Papadimitriou, and H. Ould-Brahim. Layer 1 Virtual Private Networks: driving forces and realization by GMPLS. *Communications Magazine, IEEE*, 43(7):60–67, 2005.
- [45] T. Takeda, I. Inoue, R. Aubin, and M. Carugi. Layer 1 Virtual Private Networks: service concepts, architecture requirements, and related advances in standardization. *Communications Magazine, IEEE*, 42(6):132–138, 2004.
- [46] T. Takeda, H. Kojima, and I. Inoue. Optical VPN Architecture and Mechanisms. In *The 9th Asia-Pacific Conference on Communications, APCC 2003*, volume 2, pages 751–755, 2003.
- [47] T. Takeda, H. Kojima, and I. Inoue. Layer 1 VPN architecture and its evaluation. In *The 10th Asia-Pacific Conference on Communications, APCC 2004*, volume 2, pages 612–616, 2004.

- [48] T. Takeda, H. Kojima, N. Matsuura, and I. Inoue. Resource allocation method for optical VPN. In *Optical Fiber Communication Conference, OFC 2004*, volume 1, Fevereiro 2004.
- [49] D. L. Truong, O. Cherkaoui, H. Elbiaze, N. Rico, and M. Aboulhamid. A policy-based approach for user controlled lightpath provisioning. In *Network Operations and Management Symposium, NOMS 2004*, volume 1, pages 859–872, Abril 2004.
- [50] F. Verdi, C. Carvalho, M. Magalhães, and E. Madeira. Policy-based Grooming in Optical Networks. In *4th Latin American Network Operations and Management Symposium, LANOMS 2005*, pages 125–136, Porto Alegre, Brasil, Agosto 2005.
- [51] F. Verdi, M. Magalhães, and E. Madeira. Policy-based admission control in GM-PLS optical networks. In *First International Conference on Broadband Networks Proceedings, BroadNets 2004*, pages 337–339, Outubro 2004.
- [52] D. C. Verma. Simplifying Network Administration Using Policy-Based Management. *Network, IEEE*, 16(2):20–26, 2002.