

UNIVERSIDADE FEDERAL FLUMINENSE
ESCOLA DE ENGENHARIA
MESTRADO EM ENGENHARIA DE TELECOMUNICAÇÕES

EUNICE DA CONCEIÇÃO VIEIRA PINTO AGUIAR

PROPOSTA DE SERVIÇOS EM REDES HÍBRIDAS

Niterói
2012

EUNICE DA CONCEIÇÃO VIEIRA PINTO AGUIAR

PROPOSTA DE SERVIÇOS EM REDES HÍBRIDAS

Dissertação apresentada ao Curso de Pós-Graduação “*Stricto Sensu*” em Engenharia de Telecomunicações da Universidade Federal Fluminense, como requisito parcial para obtenção do Grau de Mestre. Área de concentração: Sistemas de Telecomunicações.

Orientador: Prof. Dr. CARLOS ALBERTO MALCHER BASTOS

Niterói
2012

Ficha Catalográfica elaborada pela Biblioteca da Escola de Engenharia e Instituto de Computação da UFF

A282 Aguiar, Eunice da Conceição Vieira Pinto
Proposta de serviços em redes híbridas / Eunice da Conceição
Vieira Pinto Aguiar. – Niterói, RJ : [s.n.], 2012.
155 f.

Dissertação (Mestrado em Engenharia de Telecomunicações) -
Universidade Federal Fluminense, 2012.
Orientador: Carlos Alberto Malcher Bastos.

1. Sistema de telecomunicações. 2. Rede híbrida. 3. Fibra
óptica. 4. Roteamento. I. Título.

CDD 621.382

EUNICE DA CONCEIÇÃO VIEIRA PINTO AGUIAR

PROPOSTA DE SERVIÇOS EM REDES HÍBRIDAS

Dissertação apresentada ao Curso de Pós-Graduação “*Stricto Sensu*” em Engenharia de Telecomunicações da Universidade Federal Fluminense, como requisito parcial para obtenção do Grau de Mestre. Área de concentração: Sistemas de Telecomunicações.

Aprovada por:

Prof. Dr. Carlos Alberto Malcher Bastos, UFF

Prof. Dr. Luiz Cláudio Schara Magalhães, UFF

Prof. Dr. Carlos Ribeiro da Cunha, ESTACIO

Niterói
2012

A Deus e a todos aqueles que acreditaram que este dia chegaria.

Agradecimentos

Dedico este trabalho aos meus filhos Mariana e Pedro Henrique e ao meu marido Pedro, por se constituírem pessoas especiais e encorajadoras, estímulos que me impulsionaram a buscar uma vida nova a cada dia, meus agradecimentos por terem aceitado se privar de minha companhia pelos estudos, concedendo a oportunidade de me realizar ainda mais. Mais uma etapa de minha vida se encerra. É inevitável que para que eu chegasse ao fim deste ciclo contei com a ajuda de muitas pessoas.

Aos meus pais Elza e Wilson (*in memoriam*), minha irmã Elianete com sugestões nas suas revisões ao texto, professor Nelson que na minha formatura já me incentivou a continuar os estudos, meus companheiros de trabalho na empresa Embratel, Silvério, Fábio Benayon e Elmo com incentivos e ao Daniel Brêtas que tantas vezes me representou em reuniões que não pude comparecer, por estar participando das aulas no mestrado. Atualmente na empresa Oi, Paulo Aguiar e Wagner, pela sua paciência toda vez que peço para me ausentar e meu amigo Sérgio Coelho. Aos colegas de mestrado Bruno, Marcus, Daniel, Leandro, Francisco, Felipe e Eloy por compartilharmos diversos momentos no início dos estudos. Na fase da elaboração do projeto ao Luiz Eduardo e Ana Elisa.

Aos meus queridos Professores Schara e Débora que sempre estiveram presentes e dispostos a me ajudar, obrigada pelo carinho, dedicação e entusiasmo demonstrado ao longo do curso. Em especial ao meu Professor Orientador Malcher pelo incentivo e simpatia no auxílio às atividades e discussões, sobre o andamento e normatização desta Dissertação. Ao Professor Carlos Ribeiro pelo empenho e ajuda nas revisões e a RNP, com a oportunidade de participar das suas pesquisas de evolução de redes, com as contribuições do CPQD e de diversas Universidades Federais.

Enfim, muito obrigado a todos aqueles que direta ou indiretamente contribuíram para que este dia chegasse. Estou eternamente agradecida. E, principalmente, a DEUS pela oportunidade e pelo privilégio que nos foram dados em compartilhar tamanha experiência e, ao frequentar este curso, perceber e atentar para a relevância de temas que não faziam parte, em profundidade, das nossas vidas.

SUMÁRIO

Resumo	xv
Abstract	xvi
1. Introdução	0
1.1. Descrição do Problema	1
1.2. Motivações	2
1.3. Estrutura da Dissertação	4
2. Tecnologias de Transporte	8
2.1. Novas Tendências	8
2.2. SDH de Nova Geração	10
2.3. Tecnologia IP	12
2.4. Metro Ethernet	14
2.4.1. Ethernet Carrier Class	16
2.5. Mpls	17
2.5.1. Mpls	24
2.5.2. Gmpls	25
2.6. Redes ASON	29
2.7. Mpls-TP	36
2.8. Redes OTN	38
2.9. Redes PTN	41
2.10. Desafios da Convergência	43
2.10.1. Tendências em Serviços e Tecnologias	50
2.10.2. Novo Modelo de Redes	56
2.11. Resumo do Capítulo	57
3. Redes Híbridas e Serviços	59
3.1. Conceitos	59
3.1.1. Vantagens Redes Híbridas	62
3.2. Redes Híbridas Acadêmicas	63
3.2.1. Redes DCN	70
3.2.2. Projeto Dragon	72
3.2.3. Rede Híbrida ASON	75
3.3. Serviços Experimentais	76
3.4. Qualidade de Serviço	78
3.5. Engenharia de Tráfego	80
3.5.1. Funcionamento da Engenharia de Tráfego	81
3.6. Proteção e Restauração	83
3.7. Evolução do Perfil de Tráfego e Rede	89
3.8. Resumo do Capítulo	91
4. Proposta de Novo Serviço	92
4.1. Aplicações	93
4.2. Detalhamento do Serviço	94
4.3. Serviço UFF	97
4.3.1. Requisitos de Topologia	97
4.3.2. Classes de Serviço	98
4.3.3. Conectividade de Rede	99
4.3.4. Atributo do Serviço	100
4.3.5. Sintaxe da Aplicação Profile Service	102
4.3.6. Primitivas do Serviço	108
4.3.7. Modalidades do Serviço	112

4.4.	A Rede.....	115
4.4.1.	Suporte ao Serviço DCN.....	115
4.4.2.	Impactos na Rede.....	117
4.5.	Cenários.....	119
4.5.1.	Cenário 1.....	120
4.5.2.	Cenário 2.....	123
4.5.3.	Cenário 3.....	126
4.5.4.	Comparação dos Cenários.....	128
4.6.	Aplicativo Desenvolvido.....	129
4.6.1.	Análise dos Serviços UFF e RNP.....	129
4.7.	Resumo do Capítulo.....	132
5.	Conclusão	133
6.	Referências Bibliográficas.....	136

ÍNDICE DE FIGURAS

FIGURA 1. TAXONOMIA ATUAL	4
FIGURA 2. ESTRUTURA DISSERTAÇÃO	6
FIGURA 3. EVOLUÇÃO DAS TECNOLOGIAS	10
FIGURA 4. ESTRUTURA SDH ITU-T	12
FIGURA 5. METRO ETHERNET	17
FIGURA 6. LABEL MPLS	19
FIGURA 7. MPLS	23
FIGURA 8. ROTEAMENTO	24
FIGURA 9. G8080 ASON	30
FIGURA 10. PROTOCOLO ASON [12]	34
FIGURA 11. ITU-T G.709 [60]	40
FIGURA 12. ESTRUTURA PTN	41
FIGURA 13. NÍVEIS DE QoS	44
FIGURA 14. GRUPOS PADRONIZAÇÃO	46
FIGURA 15. MPLS-TP	47
FIGURA 16. ATRIBUIÇÕES ARQUITETURAS	49
FIGURA 17. DWDM NG ASON	50
FIGURA 18. INDICADORES DE RECEITA (TELECO)	51
FIGURA 19. CLASSIFICAÇÃO DOS SERVIÇOS	52
FIGURA 20. NOVA TAXONOMIA	53
FIGURA 21. TABELA COMPARATIVA	55
FIGURA 22. COMPARAÇÃO DOS MODELOS DE REDES	56
FIGURA 23. NOVA ARQUITETURA	60
FIGURA 24. BACKGROUND DA REDE	62
FIGURA 25. PROJETO DRAC	66
FIGURA 26. REDES HÍBRIDAS	67
FIGURA 27. PROJETO DRAGON	72
FIGURA 28. SERVIÇOS HÍBRIDOS	77
FIGURA 29. TIPOS DE TRÁFEGOS	79
FIGURA 30. MTBF DAS REDES	89
FIGURA 31. FORMATO GERAL DE MENSAGEM	94
FIGURA 32. MODELO DA REDE	95
FIGURA 33. CLASSES DE SERVIÇOS 802.1P	99
FIGURA 34. PARÂMETROS DO SERVIÇO	100
FIGURA 35. INTERFACE DE USUÁRIO E-SCIENCE	102
FIGURA 36. FORMATO DA PDU	104
FIGURA 37. WEB SERVICE	105
FIGURA 38. VISÃO GERAL	106
FIGURA 39. MAPEAMENTO VLAN	109
FIGURA 40. MAPEAMENTO FRAME	110
FIGURA 41. SERVIÇOS DA UNI	111
FIGURA 42. SELECIONA O SERVIÇO	112
FIGURA 43. REQUISITOS DOS SERVIÇOS	113
FIGURA 44. QoS PARA A DCN	116
FIGURA 45. MODALIDADES	117
FIGURA 46. CENÁRIO 1 - 1 DOMÍNIO DCN	120
FIGURA 47. SINALIZAÇÃO CENÁRIO 1	122
FIGURA 48. CENÁRIO 2 - 3 DOMÍNIOS DCN	123
FIGURA 49. SINALIZAÇÃO CENÁRIO 2	124
FIGURA 50. CENÁRIO 3 - 1 DOMÍNIO DCN ATÉ O CLIENTE	126
FIGURA 51. SINALIZAÇÃO CENÁRIO 3	127

GLOSSÁRIO

AAA – *Authentication, Authorization and Accounting*

AF - *Assured Forwarding*

AS - *Autonomous System*

ASON - *Automatically Switched Optical Network*

ASTB – *Application Specific Topology Builder*

ATM - *Asynchronous Transfer Mode*

AU - *Administrative Unit*

AUTOBAHN - *Automated Bandwidth Allocation across Heterogeneous Networks*

BoD - *Bandwidth on Demand*

BRUW - *Bandwidth Reservation for User Work*

CE – *Customer Edge*

CFQ - *Custom Fair Queuing*

CIR – *Committed Information Rate*

CLI – *Command-line Interface*

CoS – *Class of Service*

CP - *Connection Point*

CPE – *Customer Premises Equipment*

CPQD – *Centro de Pesquisa e Desenvolvimento em Telecomunicações*

CR-LDP – *Constraint-Based Routing LDP*

CS - *Client System*

CSA – *Client System Agent*

CSPF – *Constraint-Based Shortest Path First*

CTP – *Connection Termination*

DC - *Domain Controller*

DCN – *Dynamic Circuit Network*

DM - *Domain Manager*

DRAC - *Dynamic Resource Allocation Controller*

DRAGON – *Dynamic Resource Allocation via GMPLS Optical Networks*

DSCP – *DiffServ Code Point*

DWDM – *Dense Wavelength Division Multiplexing*

EF - *Expedited Forwarding*

EGP – *Exterior Gateway Protocol*

E-NNI - *External Network-Network Interface*

EPL - *Ethernet Private Line*
ERO - *Explicit Route*
EVC - *Ethernet Virtual Connection*
EVPL - *Ethernet Virtual Private Line*
FEC – *Forwarding Equivalence Class*
FSC - *Fiber-Switch Capable*
GFP - *Generic Framing Procedure*
GMPLS – *Generalized Multi-Protocol Label Switching*
GMU – *George Mason University*
GNS - *Grid Network Service*
GR - *Guaranteed Restoration*
GRE – *Generic Routing Encapsulation*
HTML - *Hyper Text Markup Language*
HTTP – *Hypertext Transfer Protocol*
ICMP - *Internet Control Message Protocol*
IDC - *Inter-Domain Controller*
IDM - *Inter-Domain Manager*
IEEE - *Institute of Electrical & Electronics Engineers*
IETF – *Internet Engineering Task Force*
IGMP - *Internet Group Management Protocol*
IGP – *Interior Gateway Protocol*
I-NNI - *Internal Network-Network Interface*
ION - *Interoperable On-demand Network*
IP - *Internet Protocol*
ISIS – *Intermediate System- Intermediate System*
ITU-T - *International Telecommunication Union Telecommunication*
JWT – *Joint Work Team*
LDP – *Label Distribution Protocol*
LER – *Label Edge Router*
LMP – *Link Management Protocol*
LRM – *Link Resource Manager*
LRO - *Label Request Object*
LSA – *Level Service Agreement*
LSC - *Lambda Switch Capable*

LSP – *Label Switched Path*
LSR – *Label Switching Router*
LTE - *Long Term Evolution*
MAC - *Media Access Control*
MAN - *Metropolitan Area Network*
MAX – *Mid-Atlantic Crossroads*
MEF – *Metro Ethernet Forum*
MEICAN - *Management Environment of Inter-domain Circuits for Advanced Networks*
MI - *Management Information*
MPLS – *Multi-Protocol Label Switching*
MPLS-TP – *Multi-Protocol Label Switching Transport Profile*
MSPP - *Multiservice Provisioning Platform*
MS-SPRING - *Multiplex Section Shared Protection Ring*
NARB – *Network Aware Resource Broken*
NCP - *Network Control Plane*
NE - *Network Element*
NGN – *Next Generation Networks*
NOC - *Network Operation Center*
NSI - *Network Service Interface*
OAM - *Operations, Administration and Maintenance*
OCh - *Optical Channel*
ODU - *Optical Data Unit*
OFS - *Optical Flow Switching*
OIF – *Optical Internetworking Forum*
OMS - *Optical Multiplex Section*
OPEX - *Operational Expenditure*
OPU - *Optical Channel Payload Unit*
OSCARS – *On-demand Secure Circuits and Advance Reservation System*
OSPF – *Open Shortest Path First*
OTN – *Optical Transport Network*
OTS - *Optical Transport Section*
OTU - *Optical Transport Unit*
OXC – *Optical Cross-Connect*
PBB – *Provider Backbone Bridge*

PCE - *Path Computation Element*
PDH – *Plesyochronous Digital Hierarchy*
PDU – *Protocol Data Units*
PE – *Provider Edge*
PerfSONAR - *Performance focused Service Oriented Network monitoring Architecture*
PoD - *Point of Demand*
POP - *Point of Presence*
PPP - *Point-to-point Protocol*
PRC - *Protection Restoration Combined*
PSC – *Packet Switch Capable*
PTN – *Packet Transport Network*
QoS – *Quality of Service*
RC – *Routing Controller*
RCE - *Resource Computation Element*
RED – *Random Early Detection*
RFC – *Request for Comments*
RIP – *Routing Information Protocol*
RNP – *Rede Nacional de Ensino e Pesquisa*
ROADM - *Reconfigurable Optical Add-drop Multiplexer*
RRO – *Record Route Objet*
RSVP - *Resource Reservation Protocol*
RSVP-TE - *Resource Reservation Protocol Traffic Engineering*
RWA - *Routing and Wavelength Assignment*
SBR - *Source Based Reroute*
SDH – *Synchronous Digital Hierarchy*
SLA – *Service Level Agreement*
SNC-P - *Subnetwork Connection Protection*
SNMP – *Simple Network Management Protocol*
SNP – *Sub-network Point*
SNPP – *Sub-network Point Pool*
SONET - *Synchronous Optical Networks*
SRG – *Shared Risk Group*
STM (-N) - *Synchronous Transport Module (-N)*
TCP - *Termination Connection Point*

TDM – *Time-division multiplexing*
TE – *Traffic Engineering*
ToIP – *Thelephony over IP*
ToS - *Type of Service*
TTL – *Time to Live*
TU - *Tributary Unit*
TUG - *Tributary Unit Groups*
UCLP - *Controlled Light Paths*
UDP - *User Datagram Protocol*
UMD – *University Maryland*
UML – *Unified Modeling Language*
UNI - *User Network Interface*
URL - *Uniform Resource Locator*
USC – *University of Southen California*
VC - *Virtual Circuit*
VLAN - *Virtual Local Area Network*
VLSR – *Virtual Label Switch Router*
VP - *Virtual Path*
VPN – *Virtual Private Network*
VRF – *Virtual Routing and Forwarding Tables*
WAN - *Wide Area Network*
WFQ - *Weighted Fair Queuing*
WG – *Working Group*
WSO – *Wavelength Switched Optical Network*
xDSL – *Digital Subscriber Line*
xPON – *Passive Optical Network*

RESUMO

Esta dissertação apresenta uma proposta de novo serviço a ser utilizado em redes híbridas. Apresenta também resumo das tecnologias ópticas e seus recentes estudos de evolução e padronizações, incluindo as redes híbridas e redes de pacotes. O serviço abrange tipos de velocidades e políticas de QoS com o processo de configuração e implantação mais fácil, pois oferece uma modelagem intuitiva para a interface e definição dos parâmetros de QoS. Permite o monitoramento das políticas de QoS através da coleta de estatísticas da rede. Este perfil está sendo modelado e pode ser implantado na rede da Universidade Federal Fluminense na cidade de Niterói, RJ e interligada à Rede Giga [68].

Palavras-chaves: Redes Ópticas, MPLS, Roteamento, Sinalização, RSVP-TE, QoS, Serviços, Redes Híbridas e DCN.

ABSTRACT

This dissertation presents a new service that may be used in hybrid networks. It also presents summary of optical technologies and their recent studies of evolution and standardization, including hybrid networks and packet networks. The service covers types of speeds and QoS policies with easier configuration and deployment process because it offers a intuitive modeling for the QoS parameters interface and definitions. Allows monitoring of QoS policies by collecting network statistics. This profile is being modeled and can be deployed in the Universidade Federal Fluminense's network, in Niteroi, RJ and connected to the Network Giga [68].

Keywords: Optical Networks, MPLS, Routing, Signaling, RSVP-TE, QoS, Services, Hybrid Network and DCN.

1. INTRODUÇÃO

Esta dissertação apresenta um novo serviço para aplicações *e-science* (método de obtenção de resultados científicos através de utilização de computação intensiva, com grande volume de dados) a ser implantado em redes híbridas (tecnologia que permite utilizar, na mesma infraestrutura de rede, conexões de pacotes e de circuitos), criando um sistema de planejamento e controle de interfaces, de forma que o serviço seja iniciado pelo cliente, através de trocas de sinalização entre os elementos dessa arquitetura e ainda modelando alguns cenários, requisitos e parâmetros necessários para viabilizar esta implementação.

As redes ópticas híbridas que possibilitam altas capacidades tiram proveito do protocolo GMPLS (*Generalized Multi-Protocol Label Switching* - [32]), com foco no plano de controle, podendo aproveitar melhor a rede legada. Com a criação de circuitos através de redes DCN (*Dynamic Circuit Network* [5]), o usuário final poderá solicitar a conexão, ou mesmo através de uma nova aplicação.

Este serviço é o de comunicação de dados em uma nova arquitetura de rede híbrida com comutação dinâmica. Esta comutação inicia com a solicitação do cliente, interagindo com o plano de controle automático da rede, com a grande vantagem de possibilitar o atendimento das necessidades de cada cliente e durante período especificado. Isto traz para as redes otimização dos recursos, flexibilidade e agilidade no atendimento à demanda, uma vez que seu estabelecimento é dinâmico. Dentro desta proposta, surgem novos conceitos como candidatos a estas arquiteturas, trazendo uma necessidade de avaliação, de como será a arquitetura de redes para implantar este novo serviço, se através de arquitetura com rede determinística ou estatística.

Um conceito surge como a DCN que é uma arquitetura de rede que permite o desenvolvimento de serviços para criação de circuitos ponto a ponto entre usuários finais com banda dedicada para algumas aplicações. Existem algumas plataformas em fase de desenvolvimento e/ou de maturação na comunidade acadêmica internacional que se propõem a implementar serviços DCN. Atualmente estão em estudo pela RNP (Rede Nacional de Ensino e Pesquisa) as plataformas DRAGON (*Dynamic Resource Allocation via GMPLS Optical Networks* - [44]), OSCARS (*On-demand Secure Circuits and Advance Reservation*

System - [54]) e AutoBahn (*Automated Bandwidth Allocation across Heterogeneous Networks*), apoiadas respectivamente pelas redes acadêmicas Internet2 e Geant [44].

1.1. DESCRIÇÃO DO PROBLEMA

Historicamente, a indústria das telecomunicações utilizou dois métodos para lidar com o aumento de tráfego: transportar dados em redes determinísticas (TDM - *time-division multiplexing*) e construir e operar em redes paralelas para o tráfego de dados de alta capacidade. A transmissão óptica de alta qualidade reduziu a necessidade de verificação da integridade dos pacotes em cada nó intermediário, o que reduz o atraso e permite a transmissão de voz sobre pacotes. Devido a essas mudanças e ao fato do tráfego de dados ser dominante, é melhor transportar voz sobre a rede de dados do que o contrário. Além disso, não é econômico continuar a construir e operar redes separadas para voz e dados [18].

Atualmente, para prover acessos de banda larga, visando atender aos serviços prestados com diversidade de velocidades requeridas e uma gama variada de requisitos [58], vem sendo utilizados os sistemas de fibras ópticas. Estes sistemas com o emprego de tecnologia DWDM (*Dense Wavelength Division Multiplexing* - G.694), oferecem largura de banda praticamente ilimitada. Poder-se-ia utilizar as revigoradas redes metálicas com a tecnologia xDSL (*Digital Subscriber Line* - G.992.1), que permitem altas taxas de velocidades, porém não são suficientes para o cenário de serviços que se anuncia [63].

Neste panorama temos as conceituadas redes SDH NG (*Synchronous Digital Hierarchy Next Generation* - [59]) bem padronizadas, porém, com operação de rede complexa, a tecnologia IP/SDH (*Internet Protocol*) que apresenta menor *overhead*, mas de difícil implementação de QoS (*Quality of Service*), a rede ASON (*Automatically Switched Optical Network* - [12]) com requisitos para comutação automática e plano de controle GMPLS, ainda com altos custos. A tecnologia metro *ethernet* mais simples e de menor custo, porém, ainda em padronização as questões de OAM (*Operations, Administration and Maintenance*).

Todas estas tecnologias estão à disposição e instaladas nas redes legadas das operadoras de Telecom, mas avançando as redes de pacotes e circuitos para uma geração de requisitos e protocolos de redes híbridas, temos o surgimento de ideias de forma a misturar

conceitos de circuitos em redes de pacotes com garantias a oferecer aos clientes, surgindo evoluções no caminho das redes híbridas PTN (*Packet Transport Network* - RFC 5654 [42]) com plano de controle GMPLS e ainda a nova geração da tecnologia DWDM NG ou OTN (*Optical Transport Network* - [60]), para prover altas capacidades nas redes *backbones*. Estas tecnologias, suas vantagens e desvantagens, bem como os atuais desafios a serem vencidos, sejam na padronização, implementação, investimentos e interoperabilidade são detalhadas no próximo capítulo.

Vemos então surgir, uma necessidade de serviços que demandem grande volume de dados com requisitos específicos, atrasos e perdas, que levam a um modelo de rede como solução destes problemas.

Este trabalho propõe uma interface que atende ao serviço dinâmico, iniciado pelo usuário, onde ele vai realizar seu *login* e senha e através da aplicação Profile Service, vai selecionar a modalidade que atende aos seus requisitos.

1.2. MOTIVAÇÕES

Nestas novas arquiteturas, os desafios de integrar as diversas plataformas existentes, sem perder os benefícios de cada uma, na evolução da tecnologia SDH (G.707), marca a chegada da tecnologia MSPP (*Multiservice Provisioning Platform*), como abordado por Sunan Han [14], como a integração das tecnologias síncronas SDH e DWDM como parte da construção das redes híbridas, e os desafios dos desenhos destas redes, como MSPP e ROADM (*Reconfigurable Optical Add-drop Multiplexer*), integrando essas plataformas.

Uma alternativa, recentemente considerada, é a definição de modelos nos quais arquiteturas dos protocolos e serviços podem ser dinamicamente programados ou adaptados ao longo de seu ciclo de vida. Essa adaptabilidade é uma característica desejável tanto para acomodar e absorver os avanços tecnológicos que acontecem cada vez mais rapidamente, como para permitir a criação de serviços e aplicações que apresentem requisitos específicos de QoS [14].

A QoS é especificada através de parâmetros para tratamento dos fluxos na rede, de forma que todos os elementos sejam coerentes entre si. A percepção do serviço pelo cliente e

pela prestadora de serviços pode ser através do tratamento do tráfego, com requisitos mínimos para se alcançar a qualidade desejada, que será detalhado com a abordagem do capítulo 4 [3].

Segundo Vicent Chan[6] na evolução das redes ópticas, precisamos identificar qual a melhor forma que podemos evoluir as redes de dados de próxima geração, utilizando as redes SDH existentes, eficientes e bem padronizadas ou as redes *ethernet* que apresentam custos relativamente mais baixos com maior simplicidade, ou mesmo, fazer um sistema híbrido.

Em um curto espaço de tempo, o cenário das telecomunicações no Brasil passou por profundas transformações. Assim chamada “convergência”, a competição e os interesses globais, estabelecem um contexto no qual achar caminhos que apontem novas direções é tarefa essencial. A convergência, associada à competição, afeta a importância e o peso econômico dos setores de telecomunicações. Com a redução das margens no negócio de transporte de informações, as operadoras de telecomunicações buscam compensar esta queda de receita pela expansão de sua base de clientes via *marketing* agressivo, pela ampliação da área de atuação e pela associação com provedores de serviço de valor adicionado e de conteúdo. Essa tendência é um movimento inevitável em direção ao compartilhamento dos recursos da rede pelas diferentes aplicações, com suporte ao protocolo IP [43]. Como características para essas redes de próxima geração, podemos citar:

- Flexibilidade para suportar a rápida implementação e a manutenção de novos serviços,
- Migração para uma arquitetura aberta onde vários serviços e tecnologias irão conviver,
- Backbone único para todos os tipos de serviço,
- Requisitos de confiabilidade / disponibilidade mais severos que os atuais,
- Predominância de serviços baseados em IP e demanda de serviços faixa larga no acesso.

Nesse modelo de redes de nova geração, uma das características é a estratificação em duas camadas: serviços e transportes. Com os avanços na camada de rede abre-se uma nova era na oferta dos serviços de voz, dados e imagem, utilizando uma única plataforma de rede, graças à convergência, que serão objeto de estudo no capítulo 3.

A camada de transporte é baseada em uma ou mais das seguintes tecnologias: OTN (G.709), SDH, ETH e MPLS-TP (*Multi-Protocol Label Switching Transport Profile*) que

baseados nos recentes padrões, dão melhorias da rede para *ethernet carrier class*. As redes de transporte ópticas são compostas por um conjunto de elementos ópticos interconectados sobre *links* de fibras, podendo prover transporte, multiplexação, roteamento, gerência, supervisão e sobrevivência, dos canais que transportam os sinais dos clientes, de acordo com a recomendação G.872 [62] (arquitetura das redes de transporte), e será objeto de estudo do capítulo 2.

1.3. ES TRUTURA DA DISSERTAÇÃO

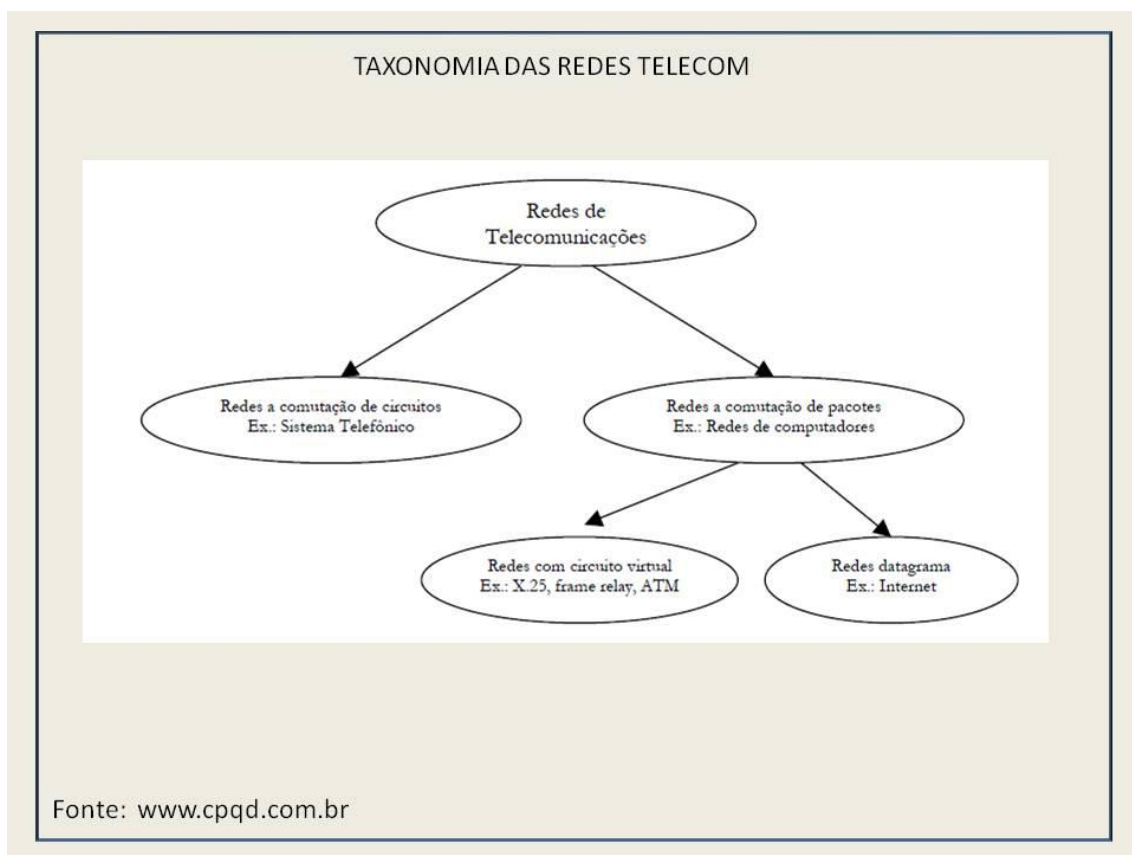


Figura 1. Taxonomia atual

A Figura 1 apresenta a taxonomia atual das redes. Suas arquiteturas são montadas em tecnologias independentes (redes de circuitos e pacotes). Segundo David Clark[9] a rede IP traz a ideia de inter-redes e as questões de roteamento, disponibilidade, gerenciamento e o argumento de levar o processamento das redes para as bordas, oferecendo aplicações distribuídas e redes de serviços.

As arquiteturas das redes NGN (*Next Generation Networks*) utilizam interfaces padronizadas e elementos de redes modulares, e todas essas tecnologias estão à disposição, dependendo das exigências dos clientes, os serviços a serem prestados como, velocidades requeridas, localização dos mesmos, etc., e poderão utilizar diferentes soluções com o emprego de diversas tecnologias quer de forma isolada ou pela combinação de algumas delas [2]. Este estudo será objeto de melhor detalhamento no capítulo 2.

Esta dissertação apresenta detalhamento das arquiteturas de rede, protocolos e funcionalidades, seus elementos integrantes e funcionamento, bem como, padrões atuais, proteção/restauração e engenharia de tráfego, através de alguns padrões IETF (*Internet Engineering Task Force* [47]) e ITU-T (*International Telecommunication Union Telecommunication* [50]), selecionados para esta dissertação e propõe serviço de dados destinados às redes híbridas, contendo detalhamento dos requisitos, especificações e implementação dentro de alguns cenários que serão discutidos no capítulo 4.

A Figura 2 contextualiza toda a abordagem desta dissertação, indicando os problemas atuais das redes em relação à evolução das redes, com o objetivo de atendimento aos serviços *e-science*. A Figura 2 também aponta alternativas para este atendimento e verificação de qual a melhor forma de atender a demanda, revendo as arquiteturas existentes ou atuando em novos padrões, e com tecnologias determinísticas ou estatísticas.

Dentro destes conceitos, algumas redes acadêmicas [44] iniciaram estudo de um plano de controle GMPLS com código aberto para que pudessem receber contribuições das demais entidades filiadas. A este plano de controle chamaram DCN, com *softwares* desenvolvidos para a criação de circuitos dinâmicos (DRAGON - *Dynamic Resource Allocation via GMPLS Optical Networks* [44]) e relacionamento entre redes autônomas (OSCARS - *On-demand Secure Circuits and Advance Reservation System* [54]).

Evoluindo as propostas acadêmicas para as redes híbridas, surge um trabalho colaborativo da RNP interligando a Redeh com a Internet2 e, ainda, a RNP adota uma solução de arquitetura híbrida para a sexta geração da sua rede *backbone*.

A UFF fez parte deste trabalho e, como contribuições, apresenta proposta quanto a serviços e suporte à serviços atendidos por rede DCN. Nesta dissertação é apresentado serviço *ethernet* com solicitação a partir do cliente, que leva em consideração questões de interoperabilidade e escalabilidade da rede legada.

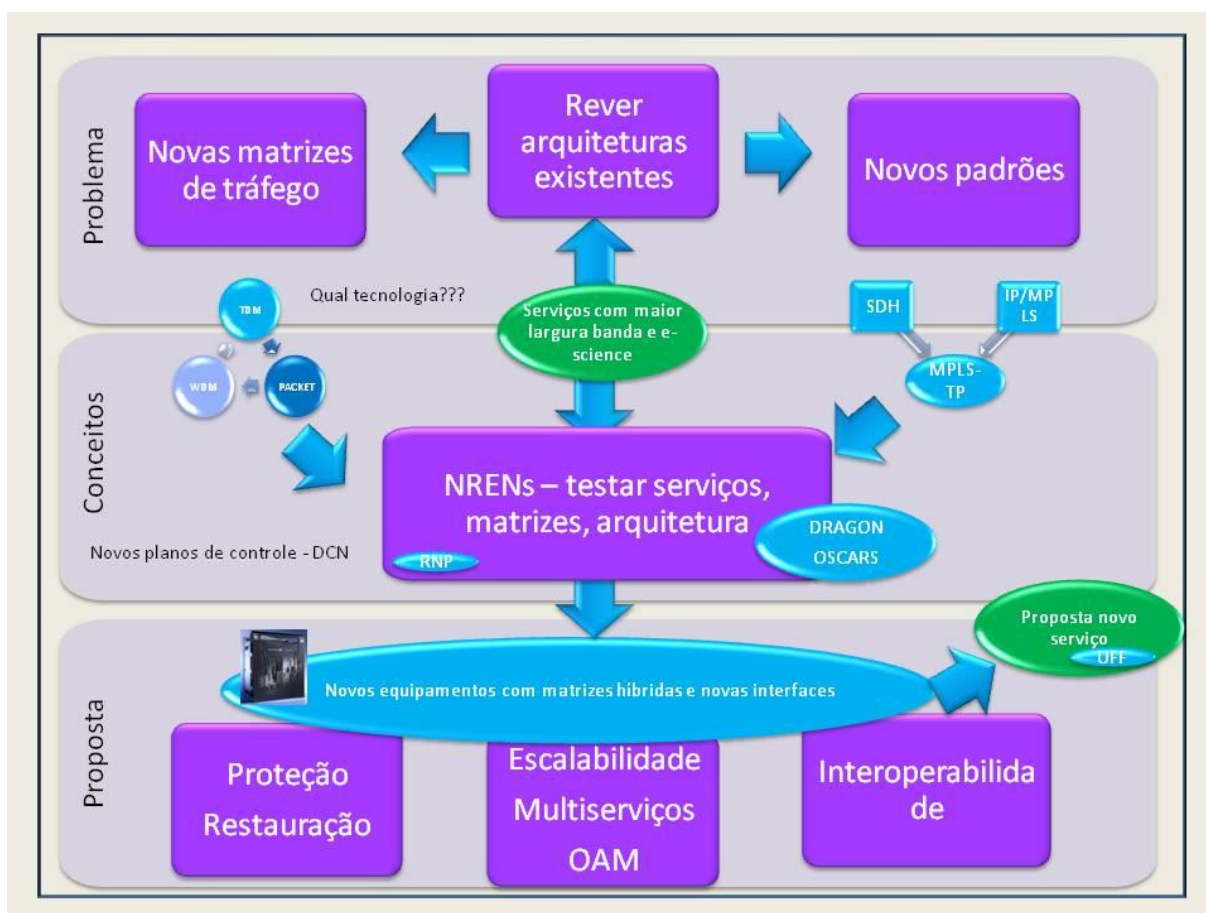


Figura 2. Estrutura Dissertação

Como parte desta dissertação, a modelagem de serviços em redes híbridas e o estudo de novas ferramentas aplicadas a rede DCN, em 2009 a UFF participa do projeto FuturaRNP, que busca a evolução da rede tendo como solução de pesquisa adotada redes híbridas. Dentro deste projeto foi implementada a Rede Cipó, uma rede de testes, superposta às redes Ipê e Giga, cujo principal objetivo é avaliar as diferentes tecnologias de redes híbridas abordadas no projeto da FuturaRNP[68], com criação de 15 *testbeds* em diversas instituições, incluindo Universidades e o CPQD (Centro de Pesquisa e Desenvolvimento em Telecomunicações).

A proposta é o estudo das soluções existentes no meio acadêmico internacional para redes híbridas, sua avaliação e adaptação para uso na RNP e testes com serviços de DCN, que contemplem também a monitoração de circuitos. A UFF no projeto FuturaRNP implementou e testou as soluções AutoBahn e DCNSS, incluindo os *softwares* OSCARS [54] e DRAGON [44]. Ambas as soluções foram testadas a cerca da facilidade de implantação e de manutenção, consistência do funcionamento e gerenciamento de alocações, entre outros critérios. Ambas as soluções, porém, ainda precisam de vários desenvolvimentos e melhorias,

em especial nos *softwares* que interagem diretamente configurando os dispositivos da rede. Falta uma infraestrutura para monitorar o estado do plano de controle e facilidades para gerenciamento, escalabilidade e segurança, são inexistentes ou muito aquém do ideal [43].

Para dar suporte técnico a esta dissertação, foi desenvolvida pesquisa sobre as redes ópticas de transporte, assunto ainda em voga e em evolução, com padrões recentes emitidos pelos institutos de padronizações: ITU-T (*International Telecommunication Union Telecommunication*), IETF (*Internet Engineering Task Force*) e OIF (*Optical Internetworking Forum*), bem como sua interação com as camadas ópticas e de pacotes, com foco no estudo das redes híbridas que oferecem serviços mistos, de comutação de pacotes e comutação de circuitos aos seus usuários, com participação no Projeto da Rede Cipó (rede híbrida experimental [68]) da RNP, conforme apresentado no *site* wiki.rnp.br.

A proposta dos novos serviços e como planejar a sua implementação é apresentada nos capítulos que seguem.

2. TECNOLOGIAS DE TRANSPORTE

2.1. NOVAS TENDÊNCIAS

Hoje estão emergindo no núcleo óptico das redes vantagens de transportar maior largura de banda e de se ter a rede com mais inteligência. Os provedores de serviço estão evoluindo para permitirem às redes manipularem múltiplos tipos de tráfego, de forma confiável e segura e ainda com níveis de QoS. Com o crescimento de serviços em tempo real e os serviços VPNs (*Virtual Private Network* - RFC 4364 [38]), aliados aos *links* de maior capacidade, as redes hoje devem acomodar novos requisitos de QoS, segurança, confiabilidade e disponibilidade e devem evoluir para os protocolos de gerenciamento de tráfego, como o MPLS-TP (RFCs 5654 / 5860) [42].

Segundo Vishnu Shukia[18] as tendências em serviços de telecomunicações nos últimos anos, foram às migrações das plataformas de voz para nova geração, como a telefonia fixa com as centrais de nova geração - NGN (ou *media gateways*) para oferta de serviços ToIP (*Thelephony over IP*), videoconferência e *home shopping*, serviços de correio eletrônico, vídeo sob demanda, TV, serviços de rede inteligente, portabilidade, além da telefonia móvel (3G e 4G).

Considerando que a evolução dessas redes tende a ser baseada em comutação de pacotes, a variação do atraso (*jitter*) e a perda de informações, precisam ser consideradas para que a qualidade do serviço ao usuário seja assegurada, portanto, uma futura rede precisará manter alto desempenho e baixo custo, para os serviços oferecidos com maior número de usuários e banda. As questões básicas para prover mecanismos de QoS ocorrem no nível de roteamento, comutação e encaminhamento dos pacotes. Existem várias soluções de alternativas mais modernas para melhorar o desempenho das redes onde a QoS é trabalhada, nas camadas 2 (dois) ou 3 (três), ou ainda, se o fluxo será identificado localmente ou fim a fim [6].

A crescente convergência das redes de serviços traz a necessidade de lidar com a QoS, que pode ser definida, como a habilidade da rede em reconhecer diferentes requisitos de serviços e diferentes aplicações de tráfego e tratamento do *delay* e *jitter*. O *delay* (RFC 2679) do datagrama marcado dentro da rede é definido como a diferença do tempo da sua entrada na

rede e o tempo de saída da rede. É chamado normalmente de latência. O *jitter* (RFC 3393) é a variação do *delay* experimentado pelos datagramas [4].

O protocolo GMPLS padronizado na RFC 3945[32], capaz de dinamicamente aprovisionar recursos e junto com as técnicas de proteção inerentes ao protocolo, torna-se fundamental nas redes ópticas. Com a criação de circuitos através de redes DCN, o usuário final poderá solicitar a conexão, ou mesmo através de uma nova aplicação e serão definidas e exemplificadas no capítulo 3.

Este capítulo descreve tecnologias de rede de transporte, relacionadas aos estudos das RFCs (*Request For Comments*) [1], arquiteturas de rede e a qualidade de serviço, que têm como objetivo, facilitar a manipulação das políticas de QoS, sejam elas para introduzir uma nova forma de desenvolvimento de políticas ou uma ferramenta de gerência de QoS, que permitam a criação de políticas e avaliação do comportamento das mesmas, ou seja, estas tecnologias vão viabilizar a rede DCN [5] que permite a implementação de serviços dinâmicos, objeto deste estudo.

As redes ópticas com comutação inteligente aliado ao protocolo GMPLS, propiciam o estabelecimento de circuitos dinâmicos, então, a rede híbrida aparece como solução neste novo cenário.

Na Figura 3 apresenta um resumo da evolução das tecnologias, que estão detalhadas nos próximos itens (2.2 a 2.9) de forma cronológica. No item 2.10 são apresentados os principais desafios das tecnologias atuais comparativamente às evoluções, seus prós e contras, além de importante visão das tendências atuais e implementação nas redes das principais operadoras no Brasil.

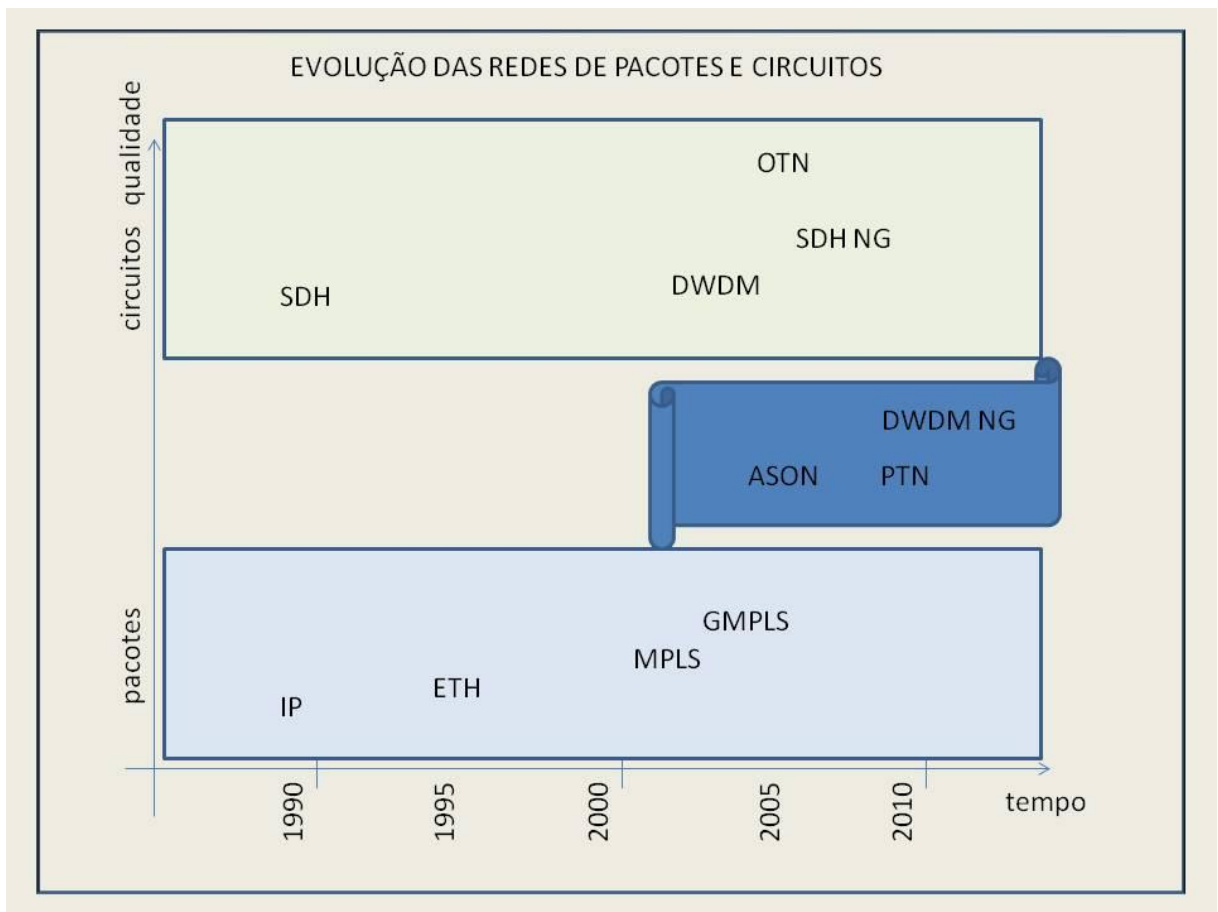


Figura 3. Evolução das Tecnologias

2.2. SDH DE NOVA GERAÇÃO

A tecnologia SDH (G.707) é hoje bem padronizada e estabelecida, com uma grande variedade de mecanismos de proteção e de interfaces (de 2 Mbps a 10 Gbps). Entretanto, o modo de transporte SDH é ineficiente para o tráfego de dados, devido à natureza estatística destes.

Existem alternativas evolutivas de uso do SDH para se adequar à natureza estatística do tráfego:

- VC (Virtual Circuit) inteiramente dedicado para o tráfego ATM (*Asynchronous Transfer Mode*) ou IP: este tipo de alternativa não permite o reuso do VC por outro tipo de tráfego (por exemplo, o tráfego determinístico – TDM).

- ATMoSDH (ATM sobre SDH): mapeamento dos VPs (*Virtual Path*) do ATM nos VCs (*Virtual Circuit*) do SDH.

-IPoATMoSDH (IP sobre ATM sobre SDH): o pacote IP é mapeado em duas células ATM de 53 octetos e depois mapeado no SDH. O resultado é uma baixa eficiência devido às perdas com os *overheads* adicionados pelas camadas.

-IPoSDH: o *payload* do SDH é preenchido com o protocolo PPP (*Point-to-point Protocol*), aumentando a eficiência, mas sem transportar tráfego TDM e células ATM.

A concatenação virtual é um mecanismo de encadeamento de entidades de transporte, para cargas que exigem capacidades maiores que os VCs especificados e permite transportar vários tipos de tráfego através da formação de agregados de tamanhos mais eficientes (por exemplo, $5 \times \text{VC-12} = 10 \text{ Mbps}$). A desvantagem deste tipo de alternativa é a alocação estática do agregado virtual.

Os *containers* virtuais padronizados podem ser utilizados para formar o quadro STM-N (*Synchronous Transport Module*). Os sinais 2 / 34 / 140 Mbit/s são mapeados nos *containers* virtuais e depois adicionados *bytes* de supervisão formando o *virtual containers* (VCs) e ponteiros formando as unidades tributárias (TUs - *Tributary Unit*). Vários TUs formam os TUGs (*Tributary Unit Groups*), que são inseridos nos *containers* virtuais de ordem superior denominado AUs (*Administrative Unit*) [59].

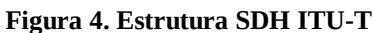
-Vantagem:

- Compatibilidade com as técnicas ATM,
- Facilidades para mistura de sinais de hierarquias diferentes em um módulo STM-1,
- Compatibilidade elétrica e óptica entre os equipamentos dos vários fornecedores,
- Padronização bem conceituada.

-Desvantagem:

- Alto custo,
- Sistema de gerência para cada tipo de equipamento,
- Operação da rede complexa.

A Figura 4 apresenta a estrutura padrão ITU-T de multiplexação do SDH, os *containers* virtuais padronizados formam o quadro STM-N, de acordo com a sua estrutura.



Os protocolos de roteamento padronizam a forma como os *gateways* trocam informações necessárias à execução dos algoritmos de roteamento.

As mensagens são associadas a cada sistema autônomo através de uma identificação na *header* (cabeçalho) da mensagem do protocolo EGP. Estas mensagens só trafegam em *gateways* (de borda) vizinhos. Dois *gateways* tornam-se vizinhos quando trocam mensagens de Aquisição de Vizinho. Após isso, verificam o estado do vizinho através da mensagem de Disponibilidade e através da mensagem Alcance identificam quais redes podem ser acessadas a partir do vizinho.

RIP (*Routing Information Protocol*) - permite a troca de informações com o algoritmo *Vector-Distance* em uma sub-rede dotada de difusão de mensagens. Um *gateway* executando RIP no modo ativo envia informações a cada 30 segundos ou quando solicitado. As mensagens contêm informações de todas as tabelas de roteamento do *gateway*. Estas informações são: endereço IP da sub-rede e a distância do *gateway* (quantidade de *gateways*). As estações e *gateways* que recebem as mensagens e atualizam sua tabela de acordo com o algoritmo *Vector-Distance*.

OSPF (*Open Shortest Path First*) - Foi desenvolvido por um grupo de trabalho da *Internet Engineering Task Force* (RFC 2328), para roteamento de grandes redes.

Utiliza o algoritmo de roteamento SPF e possui várias vantagens:

- Roteamento de acordo com o tipo de serviço,
- Balanceamento de carga entre rotas do mesmo tamanho,
- Definição de rotas específicas para máquinas e redes,
- Modularização do AS, através da criação de áreas que contém *gateways* e redes. A topologia de tais áreas é conhecida apenas nesta área,
- Definição de uma topologia de rede virtual que abstrai detalhe da rede real,
- Divulgação de mensagens recebidas de *gateways* exteriores.

Quando *gateways* OSPF são inicializados, eles verificam junto com os *gateways* vizinhos, quem será o *gateway* mestre. O *gateway* mestre será encarregado da notificação de informações de roteamento para todos os *gateways* da sub-rede. Como apenas o *gateway* mestre envia informações, o tráfego é reduzido consideravelmente [25].

Comparando com outras tecnologias temos:

- Em relação à sobreposição IP/ATM:

-Vantagem:

- Conectividade,
- Gerência de banda,
- QoS,
- Flexibilidade de reconfiguração,

- Tolerância a falhas,
- Controle de fluxo.

-Desvantagem:

- Gerenciamento de duas redes distintas,
- Número de PVs na rede,
- Roteamento IP e tráfego ATM,
- *Overhead* na ordem de 15%,
- Baixa escalabilidade.

- Em relação ao transporte IP/SDH:

-Vantagem:

- Menor *overhead* (5%),
- Tolerância a falhas,
- Velocidades das interfaces.

-Desvantagem:

- Gerência de banda,
- QoS,
- Controle de fluxo,
- Reconfiguração dos circuitos.

2.4. METRO ETHERNET

Enquanto o SDH é excelente para transportar voz, ele é ineficiente para dados, pois oferece *time slots* individuais para *streams* individuais do tráfego, o acesso baseado em tecnologia *ethernet* é uma topologia de melhor custo benefício.

São as soluções baseadas em tecnologia *ethernet* que simplificam a rede, através da redução de camadas mais complexas para aprovisionar e administrar. A solução *ethernet* em 10 Gbps pode coexistir com SDH / SONET (*Synchronous Optical Networks*), DWDM e fibra escura [51].

-Vantagem:

- Alta granularidade: de 64Kbps a 1Gbps,
- Custo menor de CPE (*Customer Premises Equipment*) para o cliente,
- Menor custo de provisionamento por eliminação das camadas de TDM / SDH,
- Agilidade de provisionamento para mudança de velocidade. A porta 10/100 *ethernet* pode ir de 64Kbps até 100Mbps sem necessidade de mudança de *hardware* e pode ser configurada por comando de supervisão remota,
- Interoperabilidade em função de serem normas simples e padronizadas,
- Simplicidade no ambiente do cliente. Não necessita especialização em ATM, *Frame Relay* e SDH,
- Interfaces ópticas padronizadas para 10Mbps, 100Mbps, 1Gbps e 10 Gbps. Há interfaces para cobre e cabo Categoria 5e para 10 e 100Mbps e Categoria 6 para 1 Gbps para cabeamentos legados,
- Facilidade para o provisionamento de banda por demanda, CIR (*Committed Information Rate*) ajustável (serviço oferecido aos clientes),
- Solução para QoS,
- Acesso final ao cliente pode ser via cobre (xDSL) e fibra.

-Desvantagem:

- Desempenho para redes locais,
- Acesso ao meio através endereçamento MAC,
- Número limite de VLAN,
- Rede não escalável.

2.4.1. ETHERNET CARRIER CLASS

A tecnologia *ethernet* tem sido largamente usada em *backbones* ou redes metropolitanas. Alguns aspectos precisam ser estudados de forma a poder participar como rede de uma operadora (*Carrier Class*), como os padrões de altas taxas em longos trechos acima de 40 km, como padronizados pelo IEEE 802.3 WG (*Working Group*), capacidades de melhoria como IEEE 802.3ah (Mac-in-Mac - *Media Access Control*), melhorias no nível dos *links* OAM [46].

A tecnologia de VLAN (*Virtual Local Area Network*) usada nos clientes que estão logicamente isoladas enquanto podem compartilhar os recursos físicos da rede, necessitando de escalabilidade no número de VLAN ID padronizada pelo IEEE 802.1ad (Q-in-Q). A escalabilidade nos *backbones* através do IEEE 802.1ah (PBB - *Provider Backbone Bridge*) que especifica B-Tag, I-Tag e C-Tag, que são os identificadores de VLAN do *backbone*, serviço e cliente, respectivamente.

Atualmente os esforços estão em prover OAM e proteção nas redes *ethernet*, como as redes SDH, permitindo aos operadores da rede detectar, localizar e verificar defeitos como facilidade e eficiência, como os padrões Y.1731 e IEEE 802.1ag e recentemente a Y.1563 (01/2009), que especifica parâmetros de *desempenho*. Com relação à comutação rápida para prover proteção à rede, a G.8031 e G.8032v2 melhoram a proteção do *Link Aggregation* e *Rapid Spanning Tree*, adicionando a capacidade da proteção em múltiplos anéis e modo não reversível [59].

As questões de QoS e controle de tráfego estão sendo estudados pelos grupos 12 e 13 do ITU-T, IEEE 802.3 e MEF (*Metro Ethernet Forum*) como o IEEE 802.1QaY (06/2009) e MEF 10.1.1 (09/2009).

Em relação às altas taxas vem sendo padronizadas as especificações IEEE 802.3av para 10G EPON (*Ethernet Passive Optical Network*) (09/2009) e a IEEE 802.3ba em desenvolvimento para 40Gbit/s e 100Gbit/s, cuja finalização dos trabalhos foi em junho de 2010.

A Figura 5 apresenta uma topologia típica de uma rede metropolitana, onde temos no *core* (*núcleo*) uma rede IP agregando o tráfego nos pontos principais da rede, podendo estar interligado ao cliente diretamente através de interfaces *fast ethernet* ou *Gigabit ethernet*, ou

com anéis coletores *ethernet* ou SDH NG, que capturam o tráfego *ethernet* dos clientes e agregam em interfaces *fast ethernet* ou *Gigabit ethernet* para então encaminhar ao *core IP*.

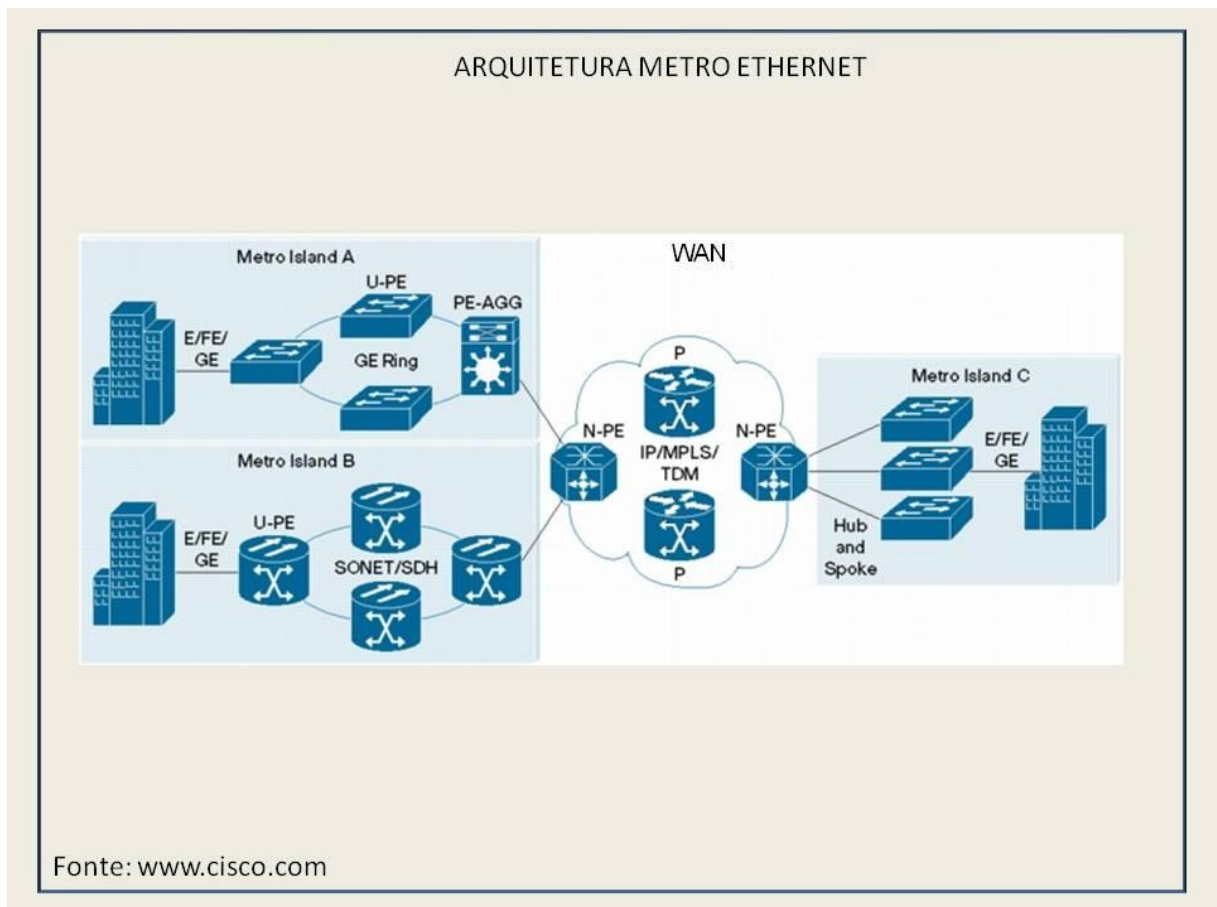


Figura 5. Metro ethernet

2.5. MPLS

O MPLS é o resultado do trabalho do IETF[27] e inspirado em experiências anteriores, como o *TAG Switching* da Cisco[65], *ARIS* da IBM[65] e *IP Switching* da Ipsilon[66]. Seu desenvolvimento iniciou em 1996 com esforços paralelos das empresas citadas acima para solucionar alguns problemas encontrados nas redes IP da época, tais como:

- Dificuldade de implementação de roteamento IP em *hardware* e implementações de *software* que limitavam a taxa de encaminhamento de pacotes nos roteadores IP.

- O encaminhamento de pacotes IP nos roteadores era ligado ao processo de roteamento IP, sem levar em conta a mudança de carga na rede, que levava ao congestionamento da rede mesmo quando alguns troncos estavam sendo subutilizados.

As redes IP eram frequentemente utilizadas sobre redes ATM, que eram extremamente dispendiosas com relação ao tamanho do cabeçalho, que era em torno de 25% a mais para todos os pacotes IP. Mas, as redes ATM tinham uma vantagem: os pacotes IP eram “forçados” para circuitos ATM particulares, sobrepondo o roteamento IP, o que aliviava o congestionamento na rede em um processo conhecido como engenharia de tráfego.

MPLS (*Multi-Protocol Label Switching*) é uma tecnologia de encaminhamento de pacotes baseada em comutação de rótulos (*label*) que funciona, basicamente, com a adição de um rótulo nos pacotes de tráfego à entrada do *backbone* e, a partir daí, todo o encaminhamento pelo *backbone* passa a ser feito com base neste rótulo. Este protocolo permite a criação de Redes Virtuais Privadas (VPN), garantindo um isolamento completo do tráfego, com a criação de tabelas de “*labels*” usadas para roteamento, exclusivas de cada VPN chamadas de tabelas VRF (*Virtual Routing and Forwarding Tables*).

Foi concebido para permitir um serviço unificado de transporte de dados para aplicações baseadas tanto em comutação de pacotes como em circuitos. Pode ser usados para transportar vários tipos de tráfego como pacotes IP, ATM, SONET ou mesmo *frames ethernet*. Esta tecnologia vem desempenhando papel importante nas redes atuais, mas ainda pode-se fazer mais com o MPLS, como uma área conhecida como engenharia de tráfego (MPLS-TE - RFC2702) [26].

Uma das vantagens do chaveamento de pacotes é a possibilidade de alterar o plano de controle do roteamento IP sem a necessidade de modificar seu algoritmo de encaminhamento.

De acordo com a RFC 3031 [27], os principais benefícios do MPLS são:

- Roteamento simplificado: O encaminhamento dos pacotes é feito baseado em *labels* curtos e de tamanho fixo que são agregados na entrada do domínio MPLS. Isto permite que os pacotes possam ser agrupados em FECs (*Forwarding Equivalence Class*) num processo conhecido como roteamento explícito.

- Roteamento explícito eficiente: O roteamento explícito é a mesma técnica que o ATM utiliza. Esta técnica permite que seja criado um caminho virtual dentro do domínio para o encaminhamento dos pacotes, evitando que estes sejam analisados em cada nó da rede.

- Suporte a engenharia de tráfego: a habilidade de definir rota dinamicamente, plano de controle dos recursos, baseado na demanda conhecida e otimização da utilização dos recursos de rede têm sido referenciados como sendo engenharia de tráfego.
- Suporte ao QoS: Além disso é possível realizar QoS com a priorização de aplicações críticas, dando um tratamento diferenciado para o tráfego entre os diferentes pontos da VPN. A QoS cria as condições necessárias para o melhor uso dos recursos da rede, permitindo também o tráfego de voz e vídeo .
- *Label* (Rótulo) - o *label* utilizado no MPLS tem tamanho fixo e curto e é utilizado como índice na tabela de encaminhamento. É agregado ao pacote IP na entrada deste no domínio MPLS e retirado quando este deixa o domínio.

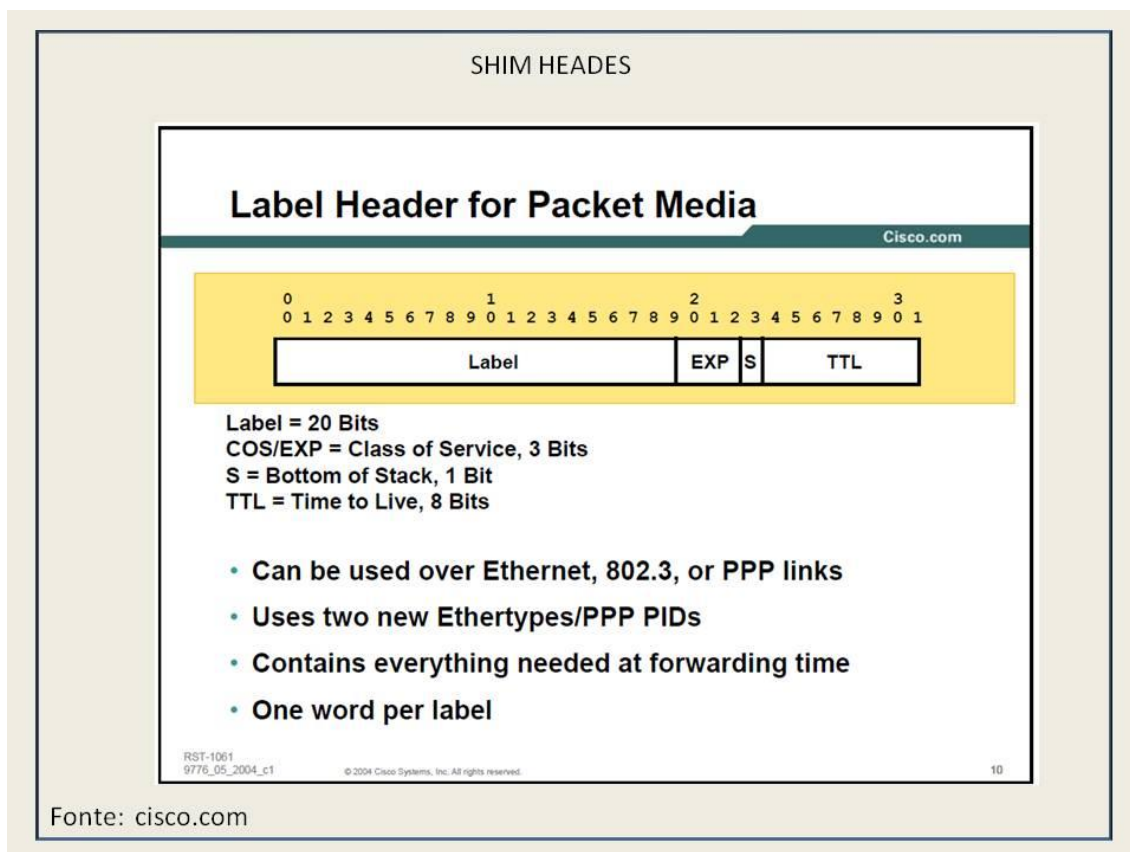


Figura 6. Label MPLS

A Figura 6 apresenta o formato do *label* MPLS - o campo CoS (*Class of Service*) é experimental e tem sido utilizado para classificar os pacotes em Classes de Serviços. O campo *stack* (S) indica se o *label* corrente está no fim da pilha. O campo TTL (*Time to Live*) tem a

mesma função que no pacote IP convencional e existe para manter a compatibilidade com ferramentas de diagnóstico de rede e para detectar *loops*.

A RFC 3036 descreve o protocolo LDP (*Label Management Protocol*) como um conjunto de procedimentos e mensagens que torna possível LSRs estabelecerem LSPs (*Label Switched Path*) na rede mapeando diretamente as informações de roteamento da camada de rede sobre os caminhos criados pela camada de enlace. O LDP associa uma FEC a cada LSP criado e estabelecendo sessões LDPs entre LSRs (*Label Switching Router*) parceiros, isto é, LSRs que trocam informações entre si e que não necessariamente são nós adjacentes na rede. Para isso o LDP define quatro tipos de mensagens que podem ser trocadas entre LSRs parceiros [28]:

- Mensagens de descobrimento: anuncia e mantém a presença de um LSR na rede.
- Mensagens de sessão: estabelece, mantém e termina sessões entre parceiros LDP.
- Mensagens de anúncio: cria, altera e finaliza mapeamento de *labels* para FECs.
- Mensagens de notificação: provê informação consultiva e sinaliza informações com erro.

É através das mensagens de descobrimento que os LSRs indicam sua presença na rede. Periodicamente LSRs enviam mensagens de “*hello*” para todos os roteadores na sub-rede. Visto que estas mensagens não necessitam de um serviço de entrega confiável o protocolo UDP (*User Datagram Protocol*) é utilizado para este fim. No entanto, para a correta operação do LDP, é necessário entrega de mensagem confiável e ordenada para os outros três tipos de mensagens, por isso, utiliza-se o protocolo TCP (*Transmission Control Protocol*).

O protocolo RSVP[24] é similar ao ICMP (*Internet Control Message Protocol* - RFC 792) ou IGMP (*Internet Group Management Protocol* - RFC 1112), padronizado pela RFC 2205. Ele permite que os nós da rede recebam informações para caracterizar fluxos de dados, definir caminhos e características de QoS para esses fluxos ao longo desses caminhos. O protocolo RSVP não é um protocolo de roteamento, ele depende de outros protocolos para execução dessas funções [24].

Este protocolo é outro esquema de distribuição de *labels* para transportes LSPs. É anterior ao surgimento do MPLS e foi concebido originalmente para criar reservas de largura de banda para fluxos de tráfego individuais numa rede. Inclui mecanismos de reserva de banda a cada salto de uma rede para uma sessão fim a fim. Entretanto, a aplicação original do

IntServ do RSVP, tinha problemas com escalabilidade, pois o número de sessões atravessando uma rede provedora de serviços seria extremamente grande.

No contexto do MPLS, entretanto, o RSVP foi estendido para permitir sua utilização na criação e manutenção de LSPs e criar reservas de banda associada. Quando usado neste contexto, o número de sessões RSVP na rede é muito menor que no caso do modelo *IntServ* por conta da maneira como o tráfego é agregado dentro do LSP. Um único LSP exige apenas uma sessão, ainda que possa transportar todo o tráfego entre roteadores, contendo muitos fluxos fim a fim. O LSP sinalizado com RSVP tem a propriedade de definir um caminho que não necessariamente o definido pelo IGP (OSPF). O RSVP, na sua forma estendida, possui propriedades de roteamento explícito em que o roteador entrante (LER - *Label Edge Router*) pode especificar todo o caminho que o LSP deve seguir ou pode especificar que o LSP deve passar por nós de trânsito particulares.

A criação do LSP sinalizado com RSVP é iniciada pelo *upstream* LER que envia mensagens de caminho RSVP. O endereço de destino da mensagem de caminho é o LER *downstream*. Entretanto, a mensagem de caminho tem a opção de ativar *Router Alert* para que roteadores de trânsito possam inspecionar o conteúdo das mensagens e realizar modificações necessárias.

Camila Cintra [8] cita alguns objetos contidos nas mensagens de caminho:

-*Label Request Object*: o LRO contém pedidos de *labels* MPLS para o caminho. Como consequência, os roteadores *upstreams* e roteadores de trânsito alocam *label* para a sessão do LSP.

-*Explicit Route Object*: o ERO contém informações de endereços de nós que o LSP deve passar. Se requerido, o ERO pode conter todo o caminho que o LSP deve seguir entre roteadores de borda, os LER.

-*Record Route Object*: o RRO solicita que o caminho seguido pela mensagem de caminho seja salva. Cada roteador por onde passa a mensagem de caminho adiciona seu endereço para a lista dentro do RRO. Um roteador pode detectar redundância de rotas caso veja seu próprio endereço no RRO.

-*Sender TSpec*: permite que o roteador entrante solicite reserva de banda para o LSP em questão. Em resposta às mensagens de caminho, o roteador entrante envia mensagens de reserva (*Resv*). Observa-se, que o roteador entrante endereça a mensagem para o roteador

vizinho ao invés de endereçar diretamente à fonte. Este esquema assegura que as mensagens *Resv* seguem exatamente o caminho inverso das mensagens de caminho [20].

Os pacotes que entram na rede MPLS são marcados com um *label* (rótulo) por um roteador de borda, mais comumente chamado de E-LSR (*Edge Label Switched Router*). Esses pacotes são encaminhados ao longo de um LSP, onde cada roteador do domínio MPLS (LSR) toma as decisões de rota baseados no conteúdo do *label*. Em cada LSR que o pacote transita, o rótulo original é substituído por outro e encaminhado ao próximo roteador. O rótulo é retirado no último roteador pertencente a esse LSP. Os roteadores de borda são os responsáveis então pela atribuição e retirada dos *labels*.

Os roteadores do núcleo de uma rede MPLS simplesmente comutam os pacotes baseados no rótulo e, portanto não efetuam consulta na tabela de roteamento, agilizando assim, o processo de encaminhamento dos pacotes.

A Figura 7 apresenta o conceito da criação dos caminhos através da rede com reserva de banda. Utiliza o modelo DSCP (*Diffserv Code Point* - RFC 2474) para diferenciação das classes, que refina a engenharia de tráfego e que habilita ao operador da rede, oferecer garantia de QoS para tráfego como voz [20].

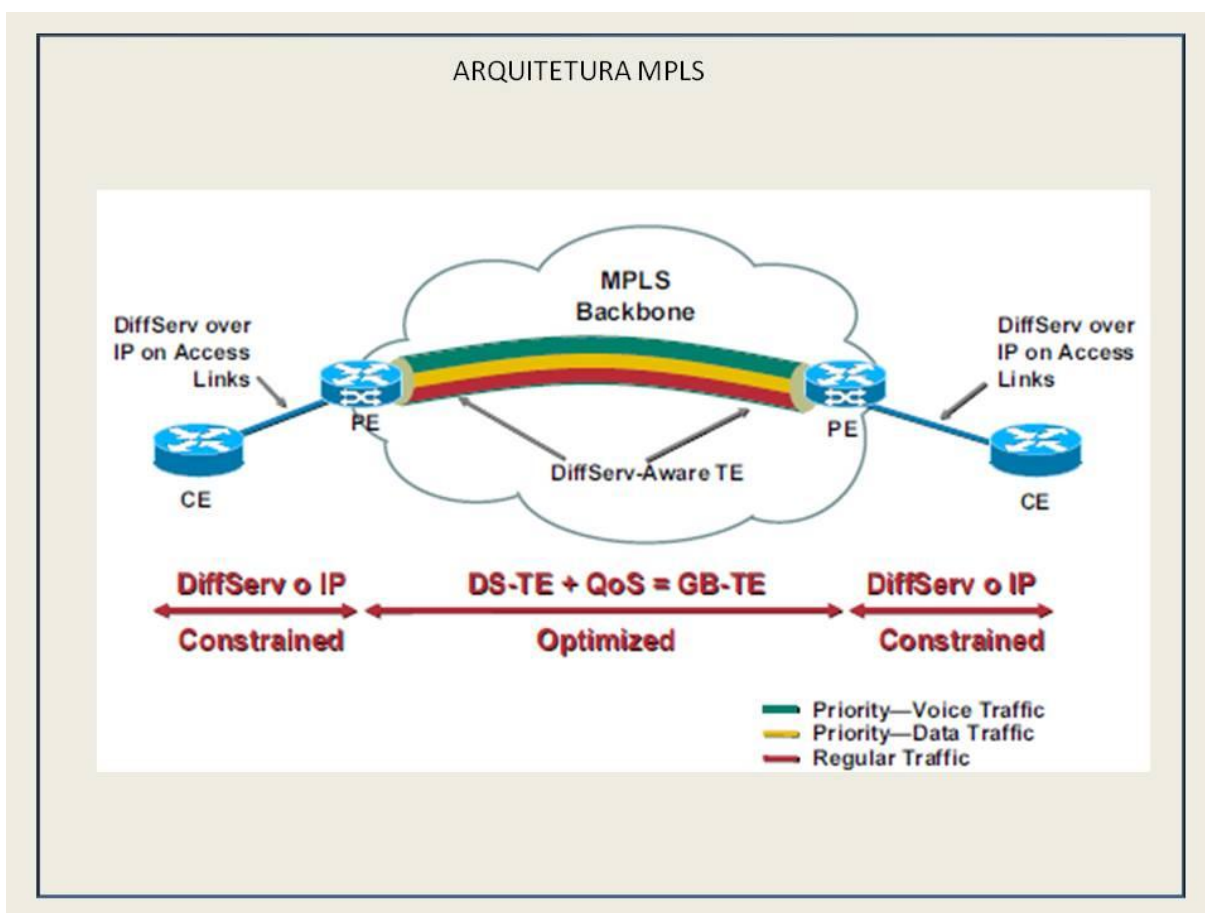


Figura 7. MPLS

O MPLS pode ser utilizado em conjunto com os algoritmos IGP para facilitar a engenharia de tráfego. O MPLS tem a informação sobre a rota específica a ser seguida por um pacote na rede. Este roteamento explícito de pacotes garante que todo o fluxo de informações com as mesmas características sigam o mesmo caminho. A administração e monitoração dos fluxos de dados permitem a avaliação correta da utilização dos recursos da rede.

Existem 3 diferenças básicas entre o roteamento convencional e o MPLS, conforme apresentados na Figura 8 e estão comparados em nível de cabeçalho, *multicast* e decisões de roteamento.

COMPARAÇÃO DOS ROTEAMENTOS		
COMPARAÇÃO	IP	MPLS
ANALISE DO CABEÇALHO IP	OCORRE A CADA NÓ	OCORRE APENAS NA BORDA
SUORTE A MULTICAST	REQUER ALGORITMOS COMPLEXOS	REQUER UM ALGORITMO
DECISÕES DE ROTEAMENTO	BASEADO NO ENDEREÇO	BASEADO EM QoS, VPN

Figura 8. Roteamento

O protocolo MPLS também tem a habilidade de avaliar informações não constantes no cabeçalho IP (por exemplo, qual o protocolo nível 4) para determinar uma rota explícita para o pacote. O administrador de rede pode desenvolver políticas de fluxo de tráfego baseadas em como e onde as informações entram na rede. Em redes tradicionais esta informação pode ser avaliada exclusivamente no ponto (roteador) de entrada. Análises adicionais permitem maior nível de controle por parte do administrador, que resulta em níveis de serviço mais previsíveis [26].

2.5.1. MPAS

Com o surgimento da tecnologia WDM foi possível multiplexar diversos comprimentos de onda em uma única fibra. Os sistemas DWDM foram inicialmente aplicados de forma a exigir operações manuais para interconexão de fibras, tornando impossível qualquer flexibilidade.

Devido à crescente importância do tráfego IP/MPLS associado com sua dinâmica de tráfego, a reconfiguração de conexão óptica tornou-se mais frequente e a realização desta

operação manualmente não era mais satisfatória. Portanto, para atender as mudanças das condições de tráfego de maneira conveniente, eliminando o enorme tempo despendido e as operações manuais possíveis causadoras de erros, os equipamentos *cross-connect* ópticos foram introduzidos para fornecer capacidades mais automatizadas e controladas remotamente. O conceito de comutação por comprimento de onda surgiu por este motivo.

Logo se tornou claro que apenas o *cross-connect* automatizado e controlado remotamente seria insuficiente para suportar aplicações de demanda tal como recuperação rápida de rede em caso de falhas. Neste caso, um plano de controle distribuído era necessário. O IETF iniciou a padronização deste plano de controle no final dos anos 90. Baseado na ideia de que um comprimento de onda pode ser considerado equivalente a um rótulo, o conjunto de protocolos MPLS foi estendido para comutação por comprimento de onda e assim surgiu o MP λ S (RFC 3471).

Segundo Furtado [11] enquanto no MPLS um par interface/rótulo de entrada é associado a um par interface/rótulo de saída, no MP λ S um par fibra/comprimento-de-onda de entrada é mapeado a um par fibra/comprimento-de-onda de saída. Depois que o trabalho de definição do MP λ S foi inicializado, percebeu-se que o paradigma *label switching* poderia não apenas ser aplicado a pacotes/células/frames e a comprimentos de onda, mas este poderia ser generalizado para *slots* de tempo, fibras, etc. Como consequência desta generalização, o MP λ S foi aprimorado gerando assim o MPLS generalizado (GMPLS – *Generalized MPLS*) [23].

A integração do IP/WDM (RFC 4209) pode então ser realizada, através da comutação de rótulos GMPLS, sendo a comutação e o roteamento realizados por meio das informações dos rótulos ópticos compostos pelos comprimentos de onda [36].

-Desvantagem:

- Comprometimento de recursos para contingência
- Confiabilidade
- Transparência do serviço

2.5.2. GMPLS

A arquitetura do plano de controle GMPLS é composta por diferentes blocos. Eles são baseados nos protocolos de sinalização e roteamento MPLS-TE que foram estendidos para suportar o GMPLS (RFCs 4203 e 4208), bem como no proposto protocolo de gerenciamento de enlace [34][35].

O roteamento GMPLS refere-se à disseminação de informação de alcançabilidade, topologia e recursos / capacidade através da topologia de roteamento do plano de controle. O GMPLS fornece extensões de roteamento TE adicionais às interfaces não PSC (*Packet Switch Capable*), se comparado ao roteamento TE realizado no MPLS-TE através dos protocolos OSPF-TE e ISIS-TE (*Intermediate System- Intermediate System*) [31].

Para o aprimoramento de escalabilidade, os seguintes mecanismos foram definidos: enlaces não numerados, tipo de proteção de enlace, informação de grupo de risco compartilhado (SRG - *Shared Risk Group*) e associação de enlaces (note que estes mecanismos são igualmente aplicáveis às interfaces PSC). Para utilizar os recursos de forma mais eficiente, o algoritmo de computação de caminhos considera várias restrições TE [32].

Assim, o roteamento baseado em restrições ultrapassa o algoritmo SPF tradicional. Estas restrições incluem restrições intrínsecas (ou de rede) tais como capacidades de enlace TE (*framing*, proteção, etc.) e suas disponibilidades de recursos (largura de banda, etc.) e, restrições extrínsecas tais como o nível de serviço requerido e o tempo médio de recuperação. Em resumo, roteamento baseado em restrições deve considerar as restrições intrínsecas e extrínsecas para computar o melhor caminho que as satisfazem.

Para tornar a computação de caminho baseada em restrições praticáveis, o LSR que inicia o pedido do LSP precisa de mais informações (atributos) sobre os enlaces TE do que os protocolos de roteamento intradomínio fornecem para as redes IP/MPLS. Estes atributos TE são disseminados usando os mecanismos de distribuição dos protocolos de roteamento de estado de enlace já estendidos para o MPLS-TE tais como, OSPF-TE e ISIS-TE. Assim, os algoritmos de computação e seleção de caminhos são mantidos específicos por cada fabricante, desde que os atributos TE sejam padronizados [33].

Extensões aos protocolos de roteamento do MPLS-TE são fornecidas para uniformemente codificarem e transportarem informações de enlace TE. Esta informação é então utilizada pelos algoritmos de computação de caminhos para estabelecerem rotas explícitas através do processo de sinalização. Rotas explícitas garantem que as restrições TE e

o nível de serviços requeridos (como parte do pedido de conexão) sejam satisfeitas no estabelecimento do LSP ao longo da rede.

O GMPLS estende dois protocolos definidos pelo MPLS-TE, o RSVP-TE e CR-LDP (*Constraint-Based Routing LDP* - RFC 3213) [29], e em alguns casos adiciona funcionalidades. Estas mudanças e adições provocam impacto em propriedades básicas dos LSP, por exemplo, na forma como os rótulos são requeridos e distribuídos, a natureza unidirecional dos LSPs, a maneira como os erros são propagados e a informação fornecida para sincronização do nó de ingresso e do nó de egresso. Estes protocolos de sinalização devem ser capazes de transportar os parâmetros LSP requeridos tais como largura de banda, tipo de sinal, proteção de enlace desejada, posição em uma multiplexação específica, etc. [40].

A sinalização GMPLS comanda a alocação de rótulo do tipo *downstream* sob demanda bem como a distribuição com controle ordenado e iniciado pelo nó de ingresso. Além disso, não há restrição quanto à estratégia de alocação de rótulos. A alocação pode ser orientada a pedido / sinalização, orientada a topologia ou até mesmo orientada a dados / tráfego (embora esta última opção seja muito comum, ela requer a definição de uma função de ativação).

Define novos blocos de construção no MPLS-TE:

- Novo formato de pedido de rótulo genérico.
- Rótulos para as interfaces PSC, LSC (*Lambda Switch Capable*) e FSC (*Fiber-Switch Capable*), conhecidos como rótulos generalizados.
- Suporte para comutação de faixa de ondas.
- Sugestão de rótulo pelo nó *upstream* para finalidade de otimização (latência).
- Restrição de rótulo pelo nó *upstream* para suportar algumas restrições ópticas.
- Estabelecimento de LSP bidirecional com resolução de contenção.
- Extensões de notificação de falha.
- Informação de proteção e foco na recuperação de enlace, mais indicação do LSP primário e secundário.
- Roteamento explícito com controle de rótulo explícito para um melhor grau de controle.

- Parâmetros de tráfego específicos por tecnologia.
- Manipulação do *status* administrativo do LSP.
- Protocolo de Gerenciamento de Enlace.

O protocolo LMP (*Link Management Protocol*), definido na RFC 4394, para gerenciamento de enlace, foi desenvolvido como parte do conjunto de protocolos GMPLS para gerenciar enlaces TE e permitir que, elementos de rede se comuniquem de forma adequada. Este protocolo permite que os nós descubram a identidade e informação detalhada de interfaces dos nós vizinhos e mantenham esta informação atualizada, assim que as configurações ou propriedades dos enlaces TE sejam modificadas [39].

O LMP consiste de quatro funções suportadas pelos controladores de canal de controle e enlace TE, as duas primeiras funções são obrigatórias e as demais são opcionais:

- Gerenciamento de canal de controle - é usado para estabelecer, configurar e manter a conectividade entre nós adjacentes (mais precisamente, entre controladores GMPLS) através da topologia de rede de controle.

- Correlação de propriedade de enlace - é usada para agregar múltiplos enlaces de dados em um único enlace TE e sincronizar a propriedade local e remota do enlace TE.

- Verificação de enlace - é usada para verificar a conectividade física dos enlaces de dados e o mapeamento entre os identificadores de interface local e remota dos enlaces de dados.

- Gerenciamento de falhas - é usado para suprimir alarmes e localizar falhas de enlace na rede [39].

Satoru Okamoto[16], descreve uma nova interface *intercarrier* desenvolvida para fornecer interoperabilidade entre os planos de controles ASON[12] e GMPLS, que inclui diferentes interfaces UNI (*User Network Interface*).

-Vantagem:

- Plano único de controle,
- Utiliza os protocolos existentes,
- Melhor aproveitamento dos recursos da rede.

-Desvantagem:

- Complexidade do algoritmo RWA (*Routing and Wavelength Assignment*),
- Interoperabilidade,
- Implementação de serviços camada 1.

2.6. REDES ASON

Especifica a arquitetura e requisitos das redes de transporte com comutação automática padronizada na G.8080, como SDH (G.803) e redes ópticas de transporte (G.872). Descreve um conjunto de componentes do plano de controle que são usados para manipular os recursos da rede de transporte, a fim de proporcionar as funcionalidades de estabelecimento, manutenção e liberação das conexões. O uso dos componentes permite a separação do controle de chamada, controle da conexão e separação do encaminhamento e sinalização [12].

Os componentes são representados como entidades abstratas, com notação similar ao UML (modelo unificado) para descrever os componentes da arquitetura da rede óptica com comutação automática. O plano de controle da rede óptica com comutação automática tem por finalidade facilitar e configurar de forma rápida e eficaz as conexões, reconfigurar e modificar conexões, realizar a restauração. A arquitetura do plano de controle pode dar aos provedores de serviço o controle de suas redes, devendo ser confiável, escalável e eficaz. Deve ser suficientemente genérica para suportar diferentes tecnologias. A recomendação trata os componentes do plano de controle e sua interação entre plano de gestão e de transporte [37].

O plano de controle ASON possui diversos componentes que administra funções específicas, incluída a determinação de rota e da sinalização. A recomendação trata os componentes da arquitetura do plano de controle e sua interação entre os planos de gerência e de transporte, embora estes dois planos estejam fora desta recomendação. A G.8080 especifica o plano de controle conforme apresentado na Figura 9.

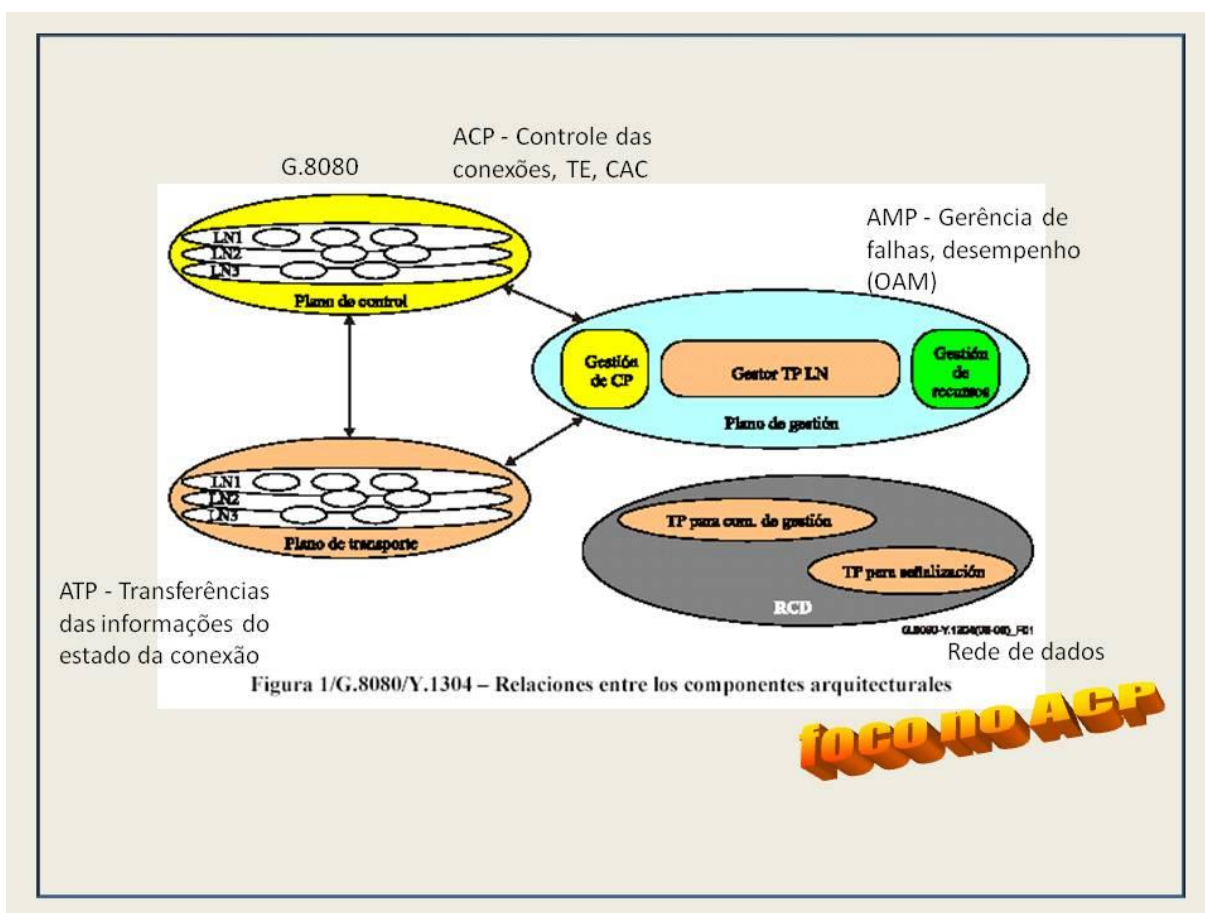


Figura 9. G8080 ASON

Controle de Chamada

O plano de controle suporta os serviços de conexão mediante o aprovisionamento automático de conexões de transporte fim a fim através de um ou mais domínios.

- Perspectiva do serviço (chamada).
- Perspectiva da conexão.

A informação relativa ao estado da conexão é detectada pelo plano de controle. O plano de controle transporta a informação relativa ao estado do enlace necessário para suportar o estabelecimento, liberação e restabelecimento da conexão. A troca de informações através dos pontos de referência é feita por múltiplas interfaces abstratas entre os componentes de controle. Uma interface física pode corresponder a várias interfaces abstratas, que podem ser multiplexadas [22].

- UNI – ponto de referência entre usuário e provedor.

- E-NNI - (*External Network-Network Interface*) ponto de referência entre domínios e provedor.
- I-NNI - (*Internal Network-Network Interface*) ponto de referência entre domínio e interface interna.

Controle de chamada – é uma associação de sinalização entre uma ou mais aplicações de usuário e a rede para controlar o estabelecimento, liberação, modificação e manutenção de conjuntos de conexões.

O controle de chamada tem que proporcionar coordenação das conexões, da seguinte forma:

- Todas as conexões devem ser encaminhadas para que possam ser supervisionadas.
- As associações de controle de chamada devem concluir antes de estabelecer as conexões.

Uma chamada tem três fases:

- Fase de estabelecimento.
- Fase de chamada ativa.
- Fase de liberação.

Controle de admissão – pode ser feita na origem da chamada verificando os parâmetros válidos, de acordo com a especificação do nível de serviço e podem ser negociados entre usuário e origem. O controle de terminação de chamada se encarrega de verificar se a parte chamada tem direito a aceitar a chamada.

Controle de conexão – se encarrega do controle global das conexões individuais, e realiza com o protocolo que estabelece e libera uma conexão. O controle de admissão da conexão é essencialmente um processo que determina se existem suficientes recursos para admitir uma conexão, se efetua enlace a enlace.

Interação entre os planos de controle, transporte e gestão:

- Gestão e transporte – funciona como um modelo de informação adequado, apresentando uma visão de recurso subjacente. Os objetos desses modelos são os recursos do transporte e interagem através das interfaces (MI - *Management Information*).

- Controle e transporte – o controlador da conexão (CC) proporciona uma interface de sinalização para controlar uma função da conexão.

- Gestão e controle – cada componente tem um conjunto de interfaces especiais que permitem supervisionar o funcionamento do componente, estabelecer políticas, interfaces (MI) do modelo funcional de transporte.

Gestão de recursos – plano de gestão, plano de controle e recursos

Recursos de Transporte e sua organização:

- Áreas de encaminhamento – conjunto de sub-redes e dos enlaces SNPP (*Sub-network Point Pool*) que interconectam [41].

- Topologia e descobrimento – as conexões de enlace que são equivalentes para fins de encaminhamento serão então agrupadas em enlaces. A informação do enlace se utiliza do LRM (*Link Resource Manager*) associados com o enlace SNPP.

- Agente de descobrimento (DA) – funciona entre o plano de transporte e proporciona a separação entre esse espaço e os nomes dos planos de controle.

- Pontos de referência – representa um conjunto de serviços UNI, I-NNI, E-NNI.

- Disponibilidade da conexão – proteção e restauração.

- Proteção – substituição de um recurso que tenha falhado por um recurso de reserva ativa.

- Restauração – substituição de um recurso que tenha falhado, por um reencaminhamento mediante uso da capacidade reserva.

- O plano de controle ASON, oferece ao um usuário chamada com CoS.

- Resiliência - capacidade do plano de controle continuar funcionando em condições de falha.

O mecanismo de proteção e restauração deve:

- Ser independente do tipo de cliente e suportar qualquer tipo de tecnologia (IP, ATM, SDH, *Ethernet*).

- Escalabilidade para suportar falha catastrófica na camada servidora, como ruptura de fibra.

- Utilizar mecanismo de sinalização robusta e eficiente, que continue funcionando mesmo após uma falha.
- Esquemas de proteção e restauração que não dependam da localização da falha.
- Proteção – supõe reencaminhamento e estabelecimento da conexão adicional.
- Restabelecimento – substituição de uma conexão em falha mediante o reencaminhamento da chamada usando a capacidade de reserva. Um domínio de reencaminhamento deve estar contido inteiramente dentro do domínio da área de controle de encaminhamento.

Quando mais de um domínio participam do reencaminhamento, os componentes de borda de cada domínio negociam a ativação dos serviços através de cada chamada. Após a comunicação ter sido estabelecida, cada um dos domínios tem conhecimento dos serviços de reencaminhamento que estão ativados para a chamada.

Durante a negociação dos serviços de reencaminhamento, os componentes de borda de um domínio trocam suas capacidades e um pedido de serviço só pode ser suportado se o serviço estiver disponível na origem e no destino e na borda de reencaminhamento.

Existem dois tipos de serviço: rígido e flexível. O rígido oferece um mecanismo de chamadas sempre em resposta a um evento de falha que em caso de falha é encaminhado, porém o *link* original é liberado antes da criação do *link* alternativo. Denomina-se “romper antes de construir”.

O mecanismo flexível para encaminhar uma conexão, estabelece uma conexão com os componentes que utilizam a principal, denominada “construir antes de romper”.

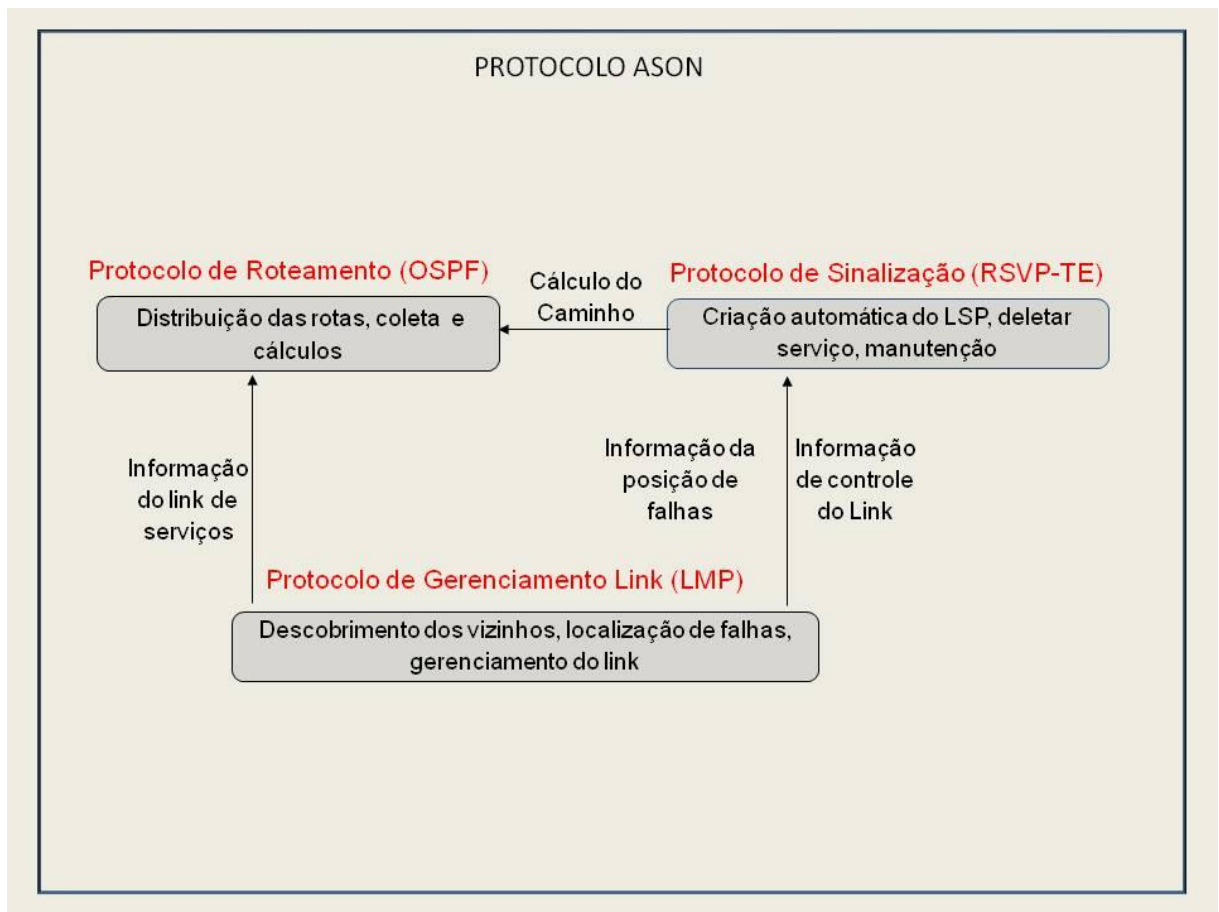


Figura 10. Protocolo ASON [12]

Conforme descrito na Figura 10, a função de encaminhamento entende a topologia em termos de enlace SNPP. Antes que se possam criar os enlaces, é preciso estabelecer a topologia de transporte e decidir as relações das conexões dos enlaces entre CPs (*Connection Point*).

Estas relações podem ser descobertas utilizando algumas técnicas, como um teste de sinal ou um *trail trace* da camada servidora ou ainda através de um sistema de gerência. As conexões do enlace que são equivalentes para fins de encaminhamento serão então agrupadas nos mesmos enlaces. Esta agregação se baseia em parâmetros, tais como: custo do enlace, retardo e qualidade ou diversidade.

A informação do enlace é também usada para configurar instâncias LRM associadas com os enlaces SNPP. Os LRMs em cada extremo do enlace têm que estabelecer uma adjacência do plano de controle que corresponda ao enlace SNPP. Uma vez concluída a validação de enlace SNPP, os LRM informam ao componente RC (*Routing Controller*) da

adjacência do enlace SNPP e as características do enlace, por exemplo, custo, desempenho e qualidade.

- Interações do controlador de chamada:

A interação entre os componentes controladores de chamada depende do tipo de chamada e do tipo da conexão.

Tipos de conexões:

- Conexões comutadas – situação na qual o controlador de chamada aceita a chamada antes que o controlador da chamada do ingresso na rede solicite.

- Conexões *soft* permanentes – o sistema de gerência envia a instrução para configurar o controle de chamada que inicia os controladores da chamada de rede.

- Processo de descobrimento do enlace:

O processo genérico de descobrimento se divide em dois instantes de tempo de espaço de nomes separados e distintos. A primeira parte ocorre internamente no espaço de nome do plano de transporte (CP e CTP - *Connection Termination*). O DA funciona internamente dentro do espaço de nome de transporte e é responsável por manter o nome do transporte da conexão de enlace. O DA consiste em um processo de descobrimento automático subjacente resolvendo cooperativamente nomes CP de transporte entre os DAs da rede. A segunda parte ocorre internamente dentro do espaço de nome do plano de controle (SNP - *Sub-Network Point*).

O gestor de recursos de enlace (LRM) mantém a informação de vinculação SNP-SNP necessária para o nome do plano de controle da conexão de enlace [12].

A ferramenta de planejamento da tecnologia ASON, permite calcular a melhor rota para cada circuito, dentro da QoS especificada e em caso de falha, permite que as rotas sejam reroteadas e compara a rota pré-calculada com os circuitos já existentes. Como *input* de entrada para a criação de novo circuito basta à definição das pontas do cliente, porém poderá ser também criado a partir de regras de estabelecimento de determinada rota, por parte do administrador, os demais circuitos serão criados automaticamente pelo ASON [30].

A ferramenta auxilia ainda sugerindo caminhos em rotas que apresentam atenuação, e vai reroteando até falhas em 5 caminhos, gerando relatórios em tempo real, podendo ao operador interagir com o sistema o tempo todo. A questão toda está na topologia da rede,

porque o tempo na montagem do circuito pode variar de 15 minutos a 24 horas, dependendo do fornecedor do equipamento.

-Vantagem:

- Aprovisionamento automático de rede: permite o reconhecimento automático do cliente;

- Suporte à escalabilidade da rede: a partir da restauração automática da conexão, o sistema será capaz de criar roteamentos do serviço em caso de falha;

- Restauração distribuída: as redes ASON propõem restauração como solução para os problemas de falha na rede ótica, permitindo a utilização da sua capacidade total.

-Desvantagem:

- Custos iniciais da rede

- Algoritmo pode demorar em realizar a convergência devido ao tamanho da rede

O ASON oferece as redes *mesh* alta disponibilidade, suportando diferentes serviços com diferentes níveis de proteção (*diamond, gold, silver*), engenharia de tráfego e ainda utiliza os protocolos existentes RSVP-TE, OSPF-TE, LMP e CSPF.

2.7. MPLS-TP

Em relação à tecnologia MPLS, o ITU (os grupos 13 e 15, desenvolvem recomendações em relação ao transporte óptico) vem unindo esforços com o IETF no sentido de interoperar os padrões existentes para as redes de transporte, o produto final deverá ser um novo *framework* conhecido como MPLS-TP, através da criação de fórum *Joint Work Team* (JWT).

Segundo Young [72] neste acordo estarão todos os requisitos para encaminhamento, OAM, proteção, gerência e protocolos de controle, cumprindo as funcionalidades das redes de transporte e garantindo interoperabilidade. As questões levantadas nos grupos de do ITU-T SG15 que estão em desenvolvimento para estabelecer os mecanismos de proteção e OAM do padrão *ethernet carrier-class*. Foram padronizados sobre OAM a Y.1710 sobre os requisitos da rede, mecanismos (Y.1711), RFC 3429, RFC 4377, RFC 4378 e RFC 4379. Foi decidido

que o grupo OAM (5/13) fosse transferido para 10/15 que deverá conduzir os estudos a partir de 2009.

A principal característica da tecnologia OTN é o transporte de qualquer sinal digital independente dos aspectos específicos do cliente. De acordo com o descrito na G.805, o limite é colocado na adaptação do canal do cliente. Temos nesta evolução em dois aspectos separados, um é o aspecto digital dos canais ópticos (ODU) com revisão da G.709, e o outro aspecto é o comprimento de onda em mídia, relacionando com a evolução para WSON (*Wavelength Switched Optical Network*) [60].

Esta nova arquitetura vai permitir que as tecnologias dos serviços e dos transportes sejam mantidas independentes pelas interfaces entre eles [72].

Ainda nesta evolução ocorre a junção dos padrões MPLS para transporte tanto do IETF como do ITU-T, de forma a criar interoperabilidade entre eles.

O desenvolvimento dos padrões do MPLS-TP (RFCs 5654 / 5860) segue os seguintes princípios [42]:

- Compatibilidade com o MPLS,
- Atender a requisitos de camada de transporte,
- Prover um conjunto mínimo de funções.

Entre outras melhorias o MPLS-TP incorporou ao seu conjunto mínimo as funções de OAM abaixo:

- Funções relativas a falhas:
 - Verificação de continuidade e conectividade: mensagens enviadas periodicamente para verificar se a conexão está normal.
 - Indicação de alarme (AIS): notificação enviada à camada cliente quando uma falha é detectada em uma camada de serviço.
 - Indicação de defeito remoto (RDI): usada para notificar a ponta remota quando ocorre uma falha local.
 - Loopback* (LB): usada para verificação bidirecional de conectividade, teste e diagnóstico bidirecional *in-service* e *out-of-service*.
 - Teste: usada para teste e diagnóstico unidirecional *in-service* e *out-of-service*.

-*Locked*: notifica que a interrupção de um serviço e sua correspondente camada ou subcamada se deve a razões administrativas, permitindo assim diferenciar uma interrupção programada de uma falha.

-*Client Signal Failure* (CSF): transfere a indicação CSF para o processo cliente específico na ponta remota.

- Funções relativas ao desempenho:

-Medição de perda de pacotes (LM): mede a taxa de perda de pacotes e *frames* tanto na ponta local como na ponta remota. Inclui métodos de medição unidirecional e bidirecional.

-Medição de retardo e variação de retardo (DM): inclui medição unidirecional e bidirecional do retardo e sua variação. A medição unidirecional requer o envio de *clock* e a sincronização de ambas as pontas, a medição bidirecional não tem este requisito.

- Outras funções de OAM (RFC 5862):

-Comutação de proteção automática (APS): provê comutação de proteção.

-Canal de gerência (MCC): provê comunicação no plano de gerência.

-Canal de sinalização (SCC): provê comunicação no plano de sinalização.

-Mensagem de sincronismo (SSM): transfere informação de sincronismo.

-Vantagem:

- Compatibilidade com o MPLS.

- Atender a requisitos de camada de transporte.

- Prover um conjunto mínimo de funções.

- OAM.

-Desvantagem:

- Finalizar padronização.

- Verificar testes de interoperabilidade.

2.8. REDES OTN

A rede óptica de transporte (OTN) foi desenvolvida com a intenção de combinar os benefícios da tecnologia SDH com a alta capacidade de expansão de banda oferecida pela tecnologia DWDM, conforme mencionamos no item 2.8. Adicionalmente as funções do SDH, o padrão do ITU-T G.709, com base na G.872, definiu a hierarquia de transporte óptico da OTN e as funcionalidades de *overhead* e ainda as estruturas do *frame*, taxa de *bits* e mapeamento dos sinais [60].

A OTN é dividida nas seguintes camadas:

- Optical Transport Section (OTS)*.
- Optical Multiplex Section (OMS)*.
- Optical Channel (OCh)*.
- Optical Transport Unit (OTU)*.
- Optical Data Unit (ODU)*.
- Optical Channel Payload Unit (OPU)*.

A tecnologia OTN adiciona ainda a correção de erro (FEC) para os elementos de rede, permitindo aos operadores diminuir o número de regeneradores necessários na rede, o que traz a vantagem da redução de custo. A OTU encapsula duas camadas, o ODU e OPU, que dão acesso à carga do SDH. Estas camadas podem ser monitoradas e possuem as seguintes taxas de linha:

- OTU1 ($255/238 \times 2.488\,320 \text{ Gb/s} \approx 2.666057143 \text{ Gb/s}$).
- OTU2 ($255/237 \times 9.953280 \text{ Gb/s} \approx 10.709225316 \text{ Gb/s}$).
- OTU3 ($255/236 \times 39.813120 \text{ Gb/s} \approx 43.018413559 \text{ Gb/s}$).
- OTU4 ($255/227 \times 99.532800 \text{ Gb/s} \approx 111.809973 \text{ Gb/s}$).

Para criar um quadro OTU a taxa do cliente é adaptada na camada ODU, que consiste em ajustar a taxa do cliente para a taxa da OPU, o *overhead* contém as informações de apoio a esta adaptação, conforme apresenta a Figura 11.

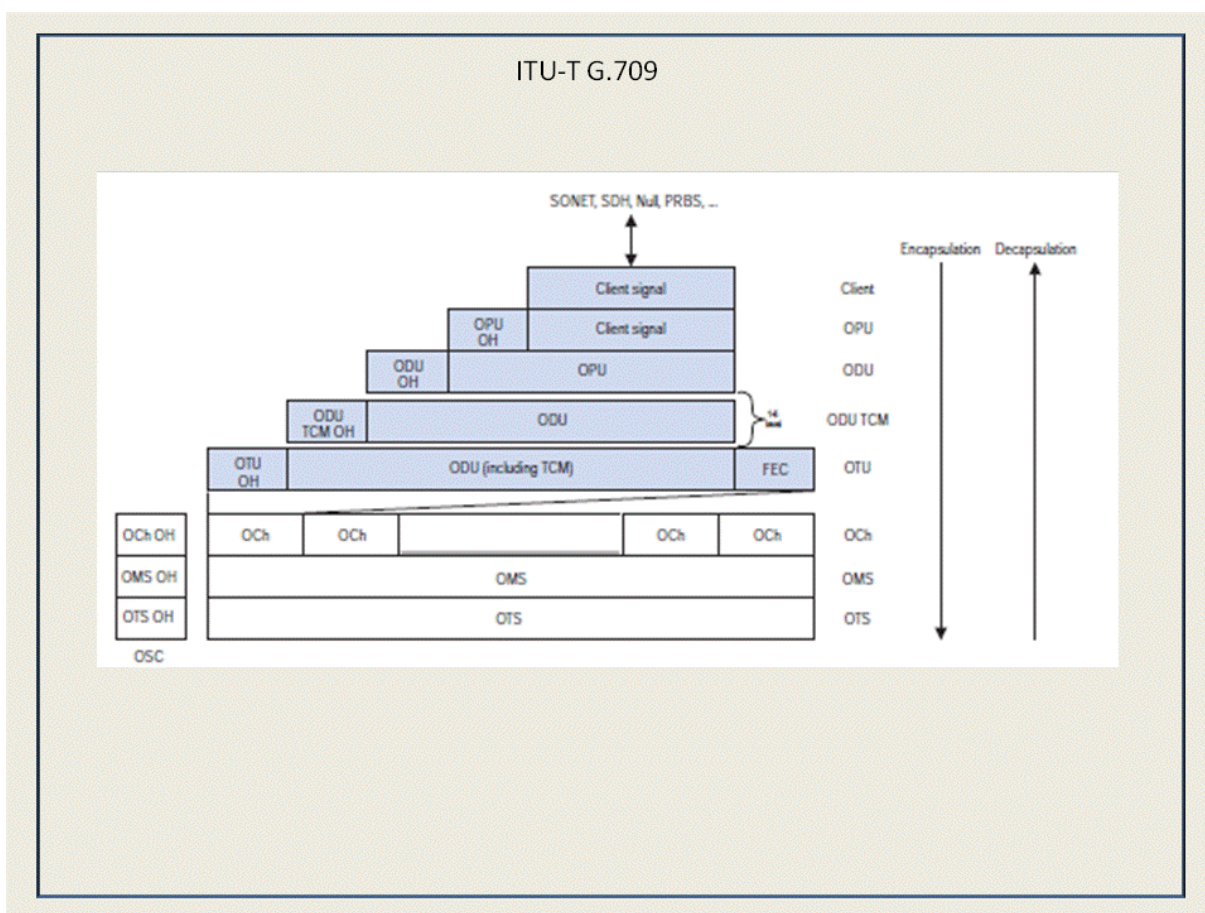


Figura 11. ITU-T G.709 [60]

Essa tecnologia continua evoluindo, mantendo os controles do SDH OAM, que permite o mapeamento mais eficiente dos sinais de dados *ethernet* e IP, com uso de *labels* MPLS-TP ou PBB-TE, orientado a conexão.

A transmissão de dados sobre esta arquitetura de rede, com monitoramento do circuito fim a fim e função de proteção, que anteriormente eram prestados pelo SDH, já estão implementados na camada WDM, com melhor eficiência e transferência de bits.

-Vantagem:

- Suporte a crescente demanda por banda.
- Suporte aos serviços 2,5 Gbps, 10 Gbps, 40 Gbps.
- Gerenciamento.
- OAM – detecção de falhas e degradação.

-Desvantagem:

- Efeitos não lineares do DWDM, ou seja, quando a energia e a matéria interagem, são efeitos inerentes a própria fibra.
- Aumento da complexidade devido às altas quantidades de canais.

2.9. REDES PTN

Conhecer as funcionalidades básicas, requisitos e especificações desta nova tecnologia, é o primeiro passo no entendimento para realização dos projetos de evolução, verificando as vantagens e desvantagens, como capacidade, disponibilidade e flexibilidade, comparativamente as tecnologias atuais.

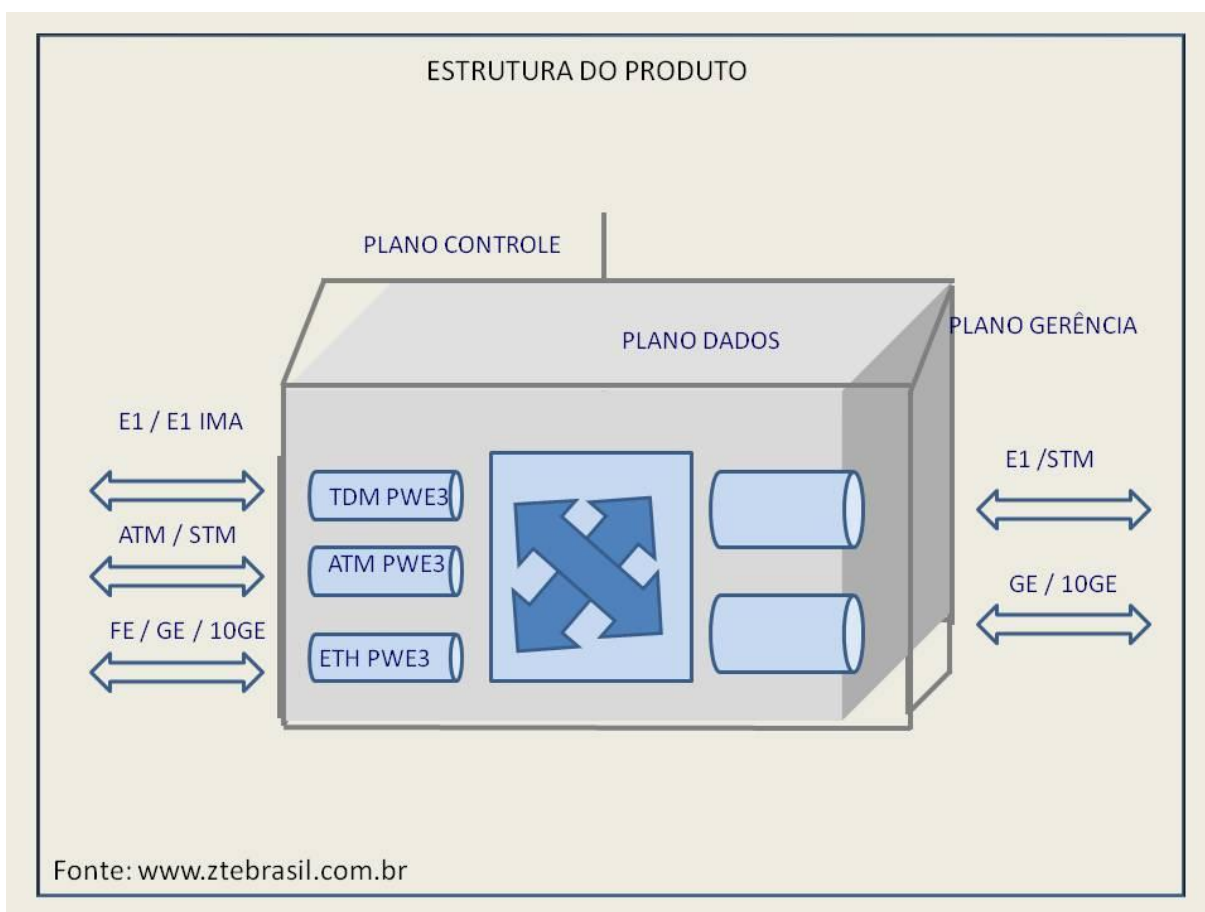


Figura 12. Estrutura PTN

A Figura 12 apresenta a estrutura da tecnologia PTN (*Packet Transport Network*) ou híbrida, os serviços são adaptados por *pseudowires* (PWE3) e em tunelamento através de MPLS-TP, antes de serem comutados por pacote. Suporta *clock* síncrono para vários tipos de

requerimentos e aplicações, podendo ser aplicados em diversos cenários, tanto nas redes de clientes, metropolitanas e *backbones*, com interfaces de alcances até 40 km.

Nesta estrutura, a placa mãe provê o tratamento dos diversos serviços em uma matriz agnóstica, responsável pelo controle do sistema, comutação e funções de sincronismo. Transmite os serviços de dados e gerência através de barramentos e componentes internos a placa. Em relação aos serviços oferecidos, poderão ser *ethernet* baseados nos padrões do MEF (E-LINE, E-LAN, E-TREE), ATM (IMA E1) ou TDM (TDM E1).

Nas funções de camada 2, vai realizar as lógicas da rede com os protocolos STP, IGMP, *Link Aggregation* e LDP. Com o MPLS-TP possibilita criar os túneis e as funções de adaptação dos serviços e ainda as funções de proteção, que poderão ser 1+1 *path protection*, 1+1 SNC, e modo *wrapping*, que é uma proteção baseada em anel. Para as funções de OAM (gerenciamento das falhas), vai realizar as verificações de continuidade e conectividade, indicação de alarmes (AIS), indicação remotas (RDI), *loopback* (conectividade bidirecional) e *lock* (trava o envio de pacotes). Em relação às funções de QoS [4] permite realizar as funções básicas, como classificação do tráfego e políticas, congestionamento, agendamento de filas, formatação do tráfego e ainda QoS dos túneis.

Esta tecnologia embora apresente interfaces para tráfego TDM, não possui placas de tributários com grandes quantidades de portas, o que significa que para o atendimento ao serviço TDM, serão necessárias várias placas, o que não otimiza o equipamento. Outra desvantagem seria em relação aos alcances até 80 km para redes de longa distancia bem inferior aos 200 km que algumas placas com *booster* são encontradas na tecnologia SDH. Outra desvantagem são os *slots* específicos para o tráfego 10GE, além do alto consumo de alguns equipamentos que chegam a 2.700 *watts*, devido às altas capacidades de processamentos das matrizes que chegam a 1,6 *terabyte*.

Como uma plataforma multisserviços, é uma rede síncrona de alta precisão, trazendo QoS ponto a ponto, OAM, mecanismos de proteção, compatibilidade com a rede instalada, etc., trazendo vantagens em relação às tecnologias atuais, tanto em capacidade como em investimento, da ordem de 50% inferior.

-Vantagem:

- Mantém as características do SDH.
- Conexão orientada.

- QoS.
 - Mecanismos de proteção.
 - OAM.
 - Escalabilidade.
 - Custos mais atrativos.
- Desvantagem:
- Ineficiente para TDM.
 - Padrões ainda em estudo.
 - Interoperabilidade.
 - Alto consumo equipamentos.

2.10. DESAFIOS DA CONVERGÊNCIA

Um dos atuais desafios nas redes de transporte é a convivência de tecnologias em desenvolvimento com a rede legada e o atendimento da demanda dos serviços tradicionais e dos novos, com uma qualidade do serviço, do mesmo nível de uma rede de circuitos, porém, em uma rede agora centrada em dados. A convergência ocorre inicialmente no serviço, que uma operadora efetivamente vende a seus clientes, porém de forma desfavorável, devido à multiplicidade de plataformas. Então, num segundo momento, surge a ideia da unificação das plataformas, com base em uma arquitetura única. Assim, do ponto de vista da rede, a convergência total é um projeto de longo prazo, em estágios a serem alcançados através do estabelecimento de padrões, tanto em arquitetura como em protocolos.

Com o avanço da utilização do protocolo IP e dos serviços baseados em *web*, boa parte das aplicações podem ser oferecidas por intermédio da *Internet* e, nestes casos, a largura de banda poderá ser oferecida com o tratamento do QoS, ou seja, um mesmo serviço poderá ser oferecido com diferentes níveis de qualidade e desempenho em função da faixa garantida e a partir de um patamar mínimo [4].

Do ponto de vista dos serviços, comparando os níveis de QoS, da aplicação e da rede, podemos destacar os protocolos RSVP, MPLS e Diffserv implementando controle em ambos os casos, o controle de filas CFQ (*Custom Fair Queuing*), WFQ (*Weighted Fair Queuing*),

RED (*Random Early Detection*), apenas aplicado à rede e o tradicional *Best Effort*, não temos aplicação de QoS nem na rede e nem na aplicação, conforme apresenta a Figura 13.

COMPARAÇÃO NÍVEIS DE QoS			
QoS	REDE	APLICAÇÃO	DESCRIÇÃO
X	X		provisionamento fim a fim
X	X	X	RSVP
X	X	X	MPLS
X	X	X	DiffServ
X	X	X	CFQ, WFQ, RED
			<i>Best effort</i>

Figura 13. Níveis de QoS

O protocolo RSVP (RFC 2205) é um protocolo de sinalização que provê mecanismos de reserva e controle para habilitar os serviços, o DiffServ provê método de classificar os serviços através de padrões de priorização com o *expedited forwarding* (EF) e o *assured forwarding* (AF).

Com relação ao mapeamento dos serviços *ethernet* sobre as redes SDH, mesmo nas redes de nova geração, a ineficiência ainda é grande. A concatenação de VCs pode ser usada para agregar o tráfego, no entanto, os esquemas de concatenação não são muito flexíveis e sua granularidade é limitada. As funções de multiplexação passam a ser executadas através de multiplexação estatística, em equipamentos ópticos integrados. Pensar na rede legada é importante, mas é preciso focar no aumento da eficiência da rede. Os roteadores atualmente estão fazendo roteamento por *hardware*, devido à evolução dos circuitos integrados, levando a um aumento de processamento e o uso de interfaces de alta velocidade, fazendo que algumas deficiências das atuais arquiteturas tenham que ser superadas.

A utilização da concatenação virtual é uma melhoria dos atuais padrões, conforme já citado, e melhora a utilização da banda. A desvantagem é que uma quantidade fixa de banda será alocada para o tráfego de pacotes, que não poderá ajustar dinamicamente o uso da banda nos diferentes VCs sob demanda, de forma estatística, ou seja, a rede ficará com alguns VCs com maior utilização enquanto outros ficarão com menor utilização.

Weiqiang Sun [17]em artigo cita os desafios atuais, para cumprimento de novos requisitos, mantendo a eficiência além de fornecer interface amigável, Observamos então a necessidade de migrar a infraestrutura de transporte de tráfego determinístico para outra, com melhor otimização e melhor custo x benefício, utilizando equipamentos que venham a ocupar espaços reduzidos e baixo consumo de energia, tornando a rede mais simples, melhor provisionamento e aumentando a competitividade, trazendo melhoria na oferta de serviços.

Dentro desta ótica, é importante acompanhar os organismos de padronização e as diretrizes de evolução, que vão trazer soluções com melhor interoperabilidade entre estas redes, permitindo que o transporte do tráfego estatístico seja otimizado nas arquiteturas de redes. Dada a importância dos protocolos, é necessário que existam padrões bem definidos. Estes são desenvolvidos pela IETF em documentos chamados de RFCs, que tendem a ser bastante técnicos e detalhados, o ITU-T com as normas (RECs) de arquiteturas de rede, o IEEE (*Institute of Electrical & Electronics Engineers* - 802.x) com as interfaces físicas e o W3C com seus padrões de conteúdos para *web*, conforme Figura 14 [19].

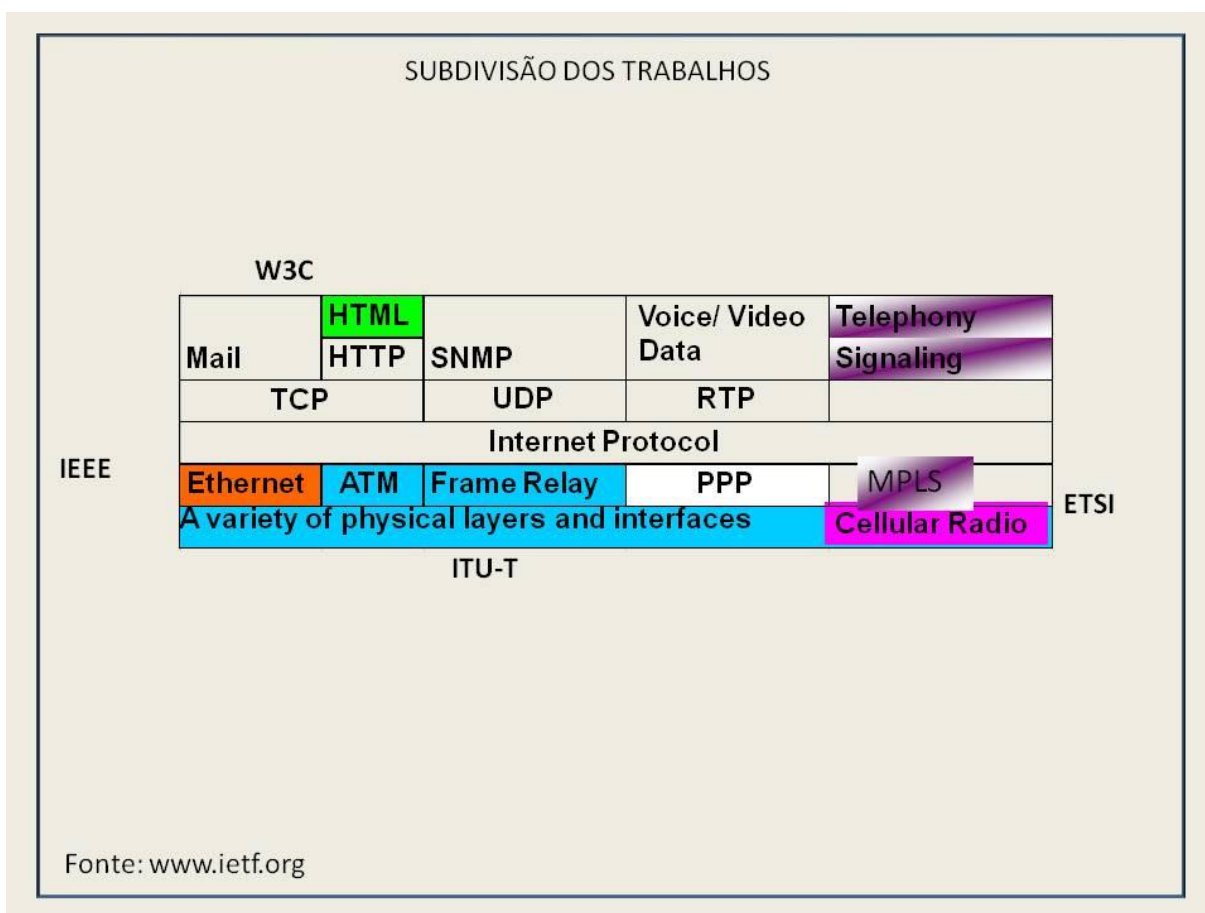


Figura 14. Grupos padronização

Soluções atuais como a *Packet Transport Network* – PTN, que se destacam das soluções já existentes de *carrier ethernet*, pelo fato de garantirem o transporte do tráfego TDM atendendo aos rígidos requisitos que requerem, permitem ainda, a sua agregação em hierarquias superiores e também o transporte e geração dos sinais de OAM, necessários à sua operação. Estas são soluções de *layer 2* que se baseiam em túneis MPLS, diferindo essencialmente nos mecanismos utilizados para garantir o transporte TDM segundo os requisitos necessários. Os padrões para esta tecnologia são definidos pelo ITU-T e IETF: o MPLS-TP (item 2.7).

Adicionalmente estão acrescentadas funcionalidades específicas para operação e manutenção: o MPLS-OAM. O ganho da multiplexação estatística que pode ser realizado usando uma rede de transporte baseada em pacotes deve traduzir uma redução de custos na rede considerando o perfil do tráfego. Essa tecnologia vem competir com os equipamentos *metro ethernet*, devido ao seu baixo custo.

Um *pseudowire* interconecta duas pontas de circuitos de um serviço através de uma rede de pacotes. Este circuito pode ser físico ou virtual e conecta o nó de cliente ao nó do provedor (CE *Customer Edge* - ao PE - *Provider Edge*), podendo ser uma porta *ethernet*, uma VLAN, um *frame relay* ou circuito virtual ATM, um E1 ou um STM-1 SDH. *Pseudowires*, que são transportados nos túneis da rede de pacotes, têm um conceito e aplicação similar aos caminhos VC4 e VC12 do SDH que são carregados dentro dos *links* STM-16 e STM-64. Os túneis na rede de pacotes que transportam os *pseudowires* são similares à seção de multiplexação do SDH. A camada do *link* (*Ethernet*, SDH, DWDM) que transportam a rede de pacotes é similar à camada de regeneração do SDH.

Para agregar sinais de informação de clientes cujas propriedades dos serviços são similares, os *pseudowires* podem ser passados de um túnel MPLS-TP para outro. *Pseudowires* que passam por vários túneis MPLS são conhecidos como *pseudowires multi-hop*, conforme arquitetura NG apresentada na Figura 15.

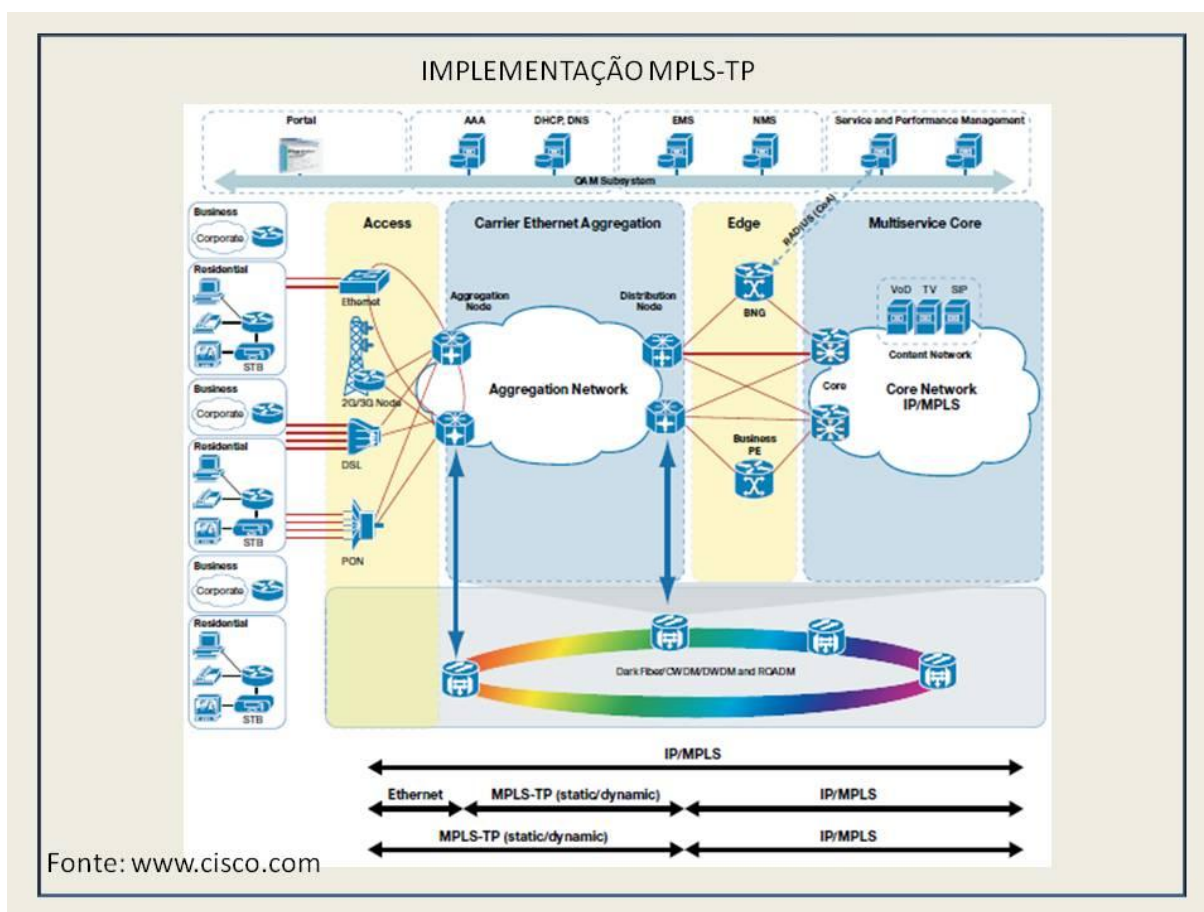


Figura 15. MPLS-TP

Este tipo de plataforma trata as questões de QoS e classificação de tráfego utilizando procedimentos de QoS hierárquico (802.1ad e VPLS) permitindo a adequação de desempenho do sistema por perfil de tráfego. Desta forma toda a engenharia de tráfego e ganhos estatísticos podem ser realizados otimizando a utilização de meios físicos e a relação de custo.

Recentemente grandes operadoras já vêm investindo em redes DWDM com matrizes híbridas e placas que agregam o tráfego com até 24 interfaces *Gigabit ethernet*, realizando a agregação em um único canal 10G, utilizando fortemente a tecnologia ASON em *backbones* nacionais e ainda em redes PTN no acesso, com matrizes agnósticas, criando novos anéis ópticos em *backbones* metropolitanos, agregando o tráfego dos equipamentos de dados, entregando esta agregação em interfaces *Gigabit ethernet* nos equipamentos *cross-connects*.

Estas matrizes são implementadas com protocolo de conversão, ou seja, elas recebem os diversos serviços tributários e converte para uma linguagem própria, realizando todo o tratamento necessário. No final do processamento reconverte e seleciona o túnel apropriado, enviando o tráfego no formato original, isto simplifica muito a construção destas novas matrizes, reduzindo o tempo de processamento. Estes equipamentos estão ganhando espaço no mercado das operadoras de forma a melhor transportar o tráfego *ethernet* e IP.

Neste panorama ainda existem diversidades de arquiteturas implementadas pelas operadoras no Brasil, enquanto umas mantêm *backbones* IP/MPLS separados da rede de transporte ASON, outras transportam todos os tipos de tráfegos na sua rede de transporte determinística, levando a aumentos das suas capacidades a cada 6 meses. Um ponto que vem sendo estudado desde 2011 é a evolução da tecnologia WDM NG ASON, trazendo *transponders* de 40G e 100G com matrizes O-O-O (óptica-óptica-óptica), de forma a migrar o tráfego estatístico para esta plataforma, deixando que a rede legada fique com reserva de capacidade. A Figura 16 apresenta as diferentes arquiteturas para transporte do tráfego IP.

ATRIBUIÇÕES DAS ARQUITETURAS		
ARQUITETURA	VANTAGENS	DESVANTAGENS
IP/ATM	CONECTIVIDADE. ENGENHARIA TRÁFEGO. QoS	ALTO OVERHEAD. DOIS PLANOS DE CONTROLE. GERENCIAMENTO DE 2 REDES DISTINTAS. ESCALABILIDADE
IP/SDH	MENOR OVERHEAD. GRANDE VAZÃO	CONECTIVIDADE. NUMERO LIMITADO DE ACESSOS WAN
IP/MPLS	ESCALABILIDADE. TE. QoS. CONECTIVIDADE	PADRONIZAÇÃO
IP/MPAS	INTERCONEXÃO DIRETA COM A CAMADA ÓPTICA	PADRONIZAÇÃO

Figura 16. Atribuições arquiteturas

Nestas novas redes WDM NG ASON (Figura 17), as características da rede ASON serão utilizadas, onde os NE (*Network Element*) poderão buscar o NE vizinho automaticamente, construir sua própria topologia e o serviço será fim a fim baseado em comprimento de onda (*Och trail*) de forma configurável pelo equipamento. O administrador da rede precisa especificar o NE de origem e de destino, a largura de banda necessária e o nível de proteção para criação de circuito. O roteamento do serviço e a *cross-conexão* dos NEs intermediários serão automaticamente criados pelo ASON.

O protocolo de roteamento seleciona a melhor rota baseando-se em 3 condições: distância de transmissão, saltos e balanceamento da banda.

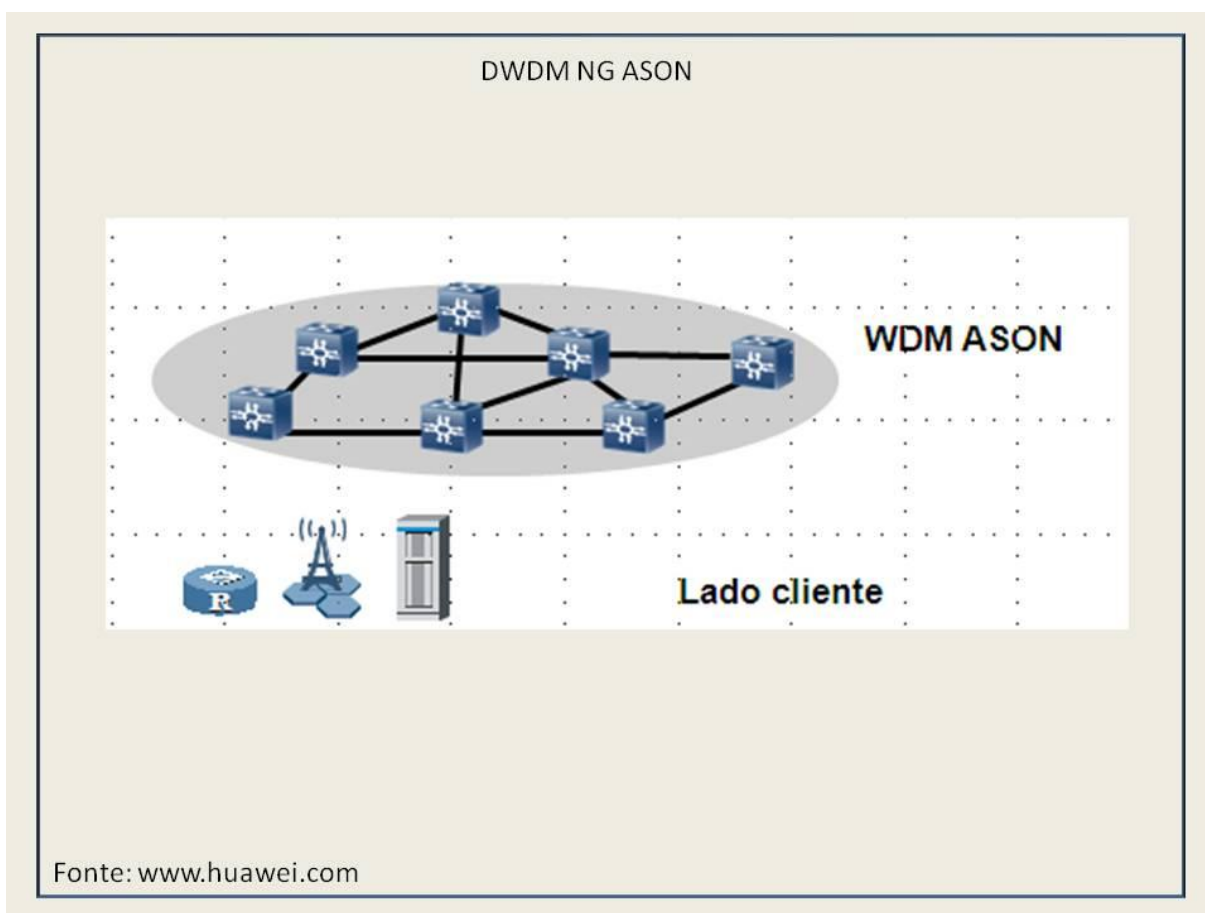


Figura 17. DWDM NG ASON

Durante o planejamento da rede, no caso em que a topologia for muito complexa, será necessário subdividir em sub-redes de acordo com o serviço. Para iniciar o planejamento, terá início com base na análise da camada física, verificando os fatores como: potência óptica, dispersão, etc. Neste caso, serão necessários algoritmos relevantes, com o cuidado de utilizar um menor número de métricas. Nesta plataforma o plano de controle será com o protocolo GMPLS e a utilização de ferramenta *off line* para projetar e verificar simulação de falhas, antes de serem aplicadas a rede, aplicando a cada comprimento de onda o nível de serviço, tipo de proteção e restauração desejada [53].

2.10.1. TENDÊNCIAS EM SERVIÇOS E TECNOLOGIAS

O serviço de voz ainda é dominante para valores de receita (Figura 18), pois significa uma fatia dos negócios que deverão garantir o *market share* para as operadoras estabelecidas.

Então as aplicações deverão estar baseadas em redes inteligentes e que tenham escalabilidade, aliado aos riscos com as competições, isto tudo, irá causar um forte efeito na escolha das tecnologias.

DESEMPENHO ECONOMICO FINANCEIRO			
Quebra da Receita Líquida			
R\$ Milhões	2010	2011	2012
Receita Líquida	29.479	27.907	28.142
Residencial	11.949	10.501	9.974
Mobilidade Pessoal	8.021	8.190	9.102
Serviços	7.917	8.154	8.612
Voz*	4.958	5.137	6.276
Uso de Rede	2.305	2.398	2.337
Dados / Valor Adicionado	654	620	-
Material de Revenda	104	36	489
Empresarial / Corporativo	8.620	8.470	8.510
Outros serviços	890	746	556

*Inclui Dados a partir de 2012.

Fonte: teleco.com.br/Operadoras/oibrt.asp

Figura 18. Indicadores de Receita (Teleco)

A Figura 19 apresenta uma taxonomia padrão para a classificação dos serviços de acordo com o tipo de comutação [64].



Figura 19. Classificação dos Serviços

Como rede/tecnologia legada, temos a camada de transporte utilizando as tecnologias de camada 1 do SDH e do DWDM, que se desenvolveram no transporte dos serviços síncronos, sendo que transporte do tráfego *ethernet* melhor se adapta à evolução do SDH e do DWDM com suas interfaces transparentes. Uma evolução que naturalmente ocorre pela combinação dos benefícios da tecnologia SDH com a alta capacidade de banda oferecida pelo DWDM, é o padrão OTN, que está se capacitando para oferecer *lambdas* de 100 Gbps no padrão DWDM NG.

Em paralelo a esta grande evolução, temos a camada de serviços utilizando as tecnologias IP e MPLS, com forte tendência em utilizar para os serviços assíncronos, uma qualidade de serviço com melhor utilização das bandas dos *links* da rede. Este paralelo levou em 2006 ao estabelecimento de novas redes baseadas em plano de controle, com a chegada do protocolo GMPLS e as redes ASON. Este foi o marco principal para as operadoras começarem a ter melhor controle das conexões e estabelecimento de circuitos automatizados.

Porém, as redes ainda apresentavam problemas de interoperabilidade entre elas. Isto levou a OIF a criar estruturas de interfaces padronizadas. Surge, assim, o padrão MPLS-TP, que acrescenta ao conceito das redes híbridas, a comutação por rótulos e QoS. As redes de nova geração são criadas onde podemos utilizar, sobre a mesma infraestrutura de rede, conexões de pacotes e de circuitos.

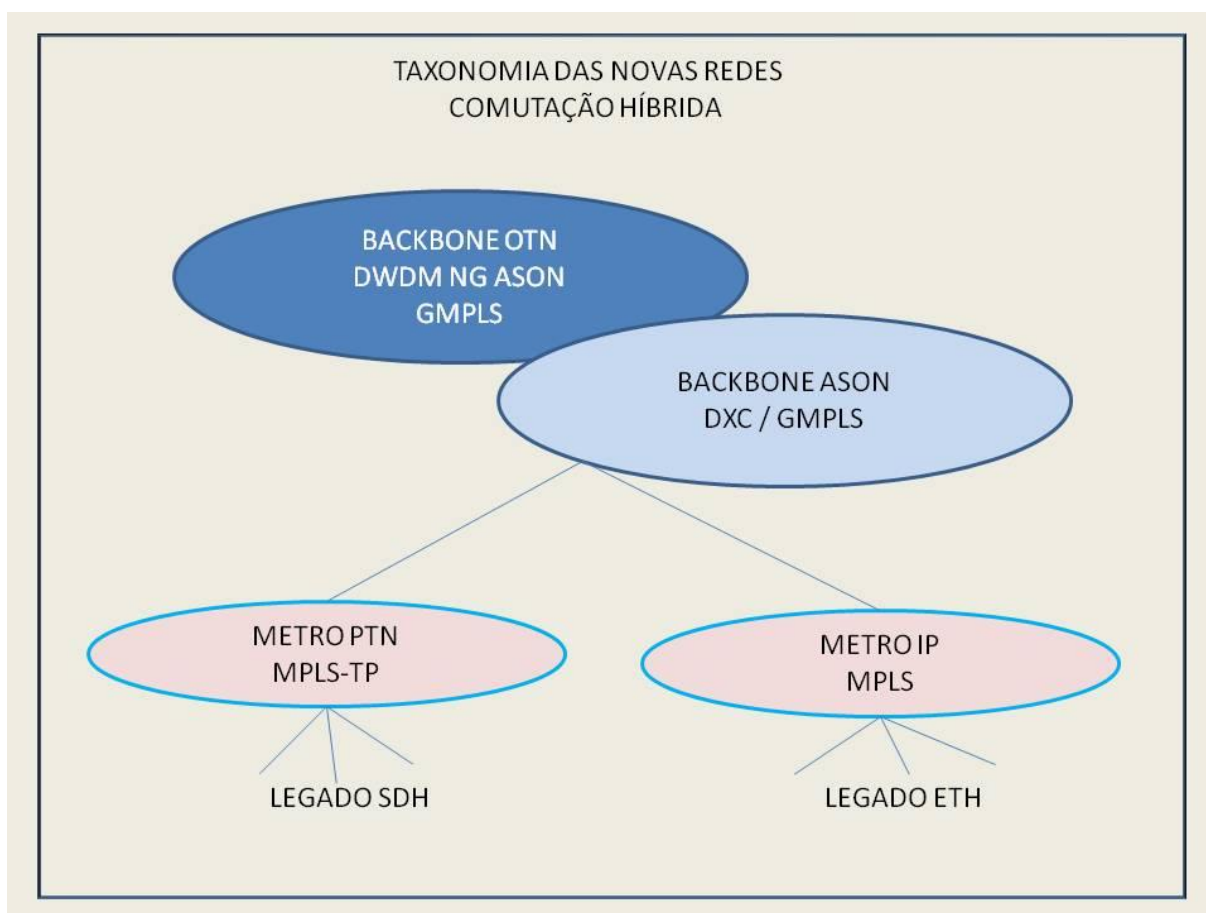


Figura 20. Nova taxonomia

A Figura 20 apresenta esta ideia de evolução, onde vemos as camadas de transporte e de serviços convergindo para padrões mais robustos e adaptados às novas demandas de serviços, ou seja, adotando as melhores tecnologias para cada tipo de segmento da rede que se deseja expandir. Por exemplo, para as redes de acesso pode-se utilizar as redes PTN (com o protocolo MPLS-TP) e para os *backbones* as redes DWDM NG (*lambdas* de 40G e 100G).

Neste novo cenário, com diferentes atores (consumidores e políticas de Governo), um dos principais desafios será a mudança do paradigma da visão de mercado pela estratégia de negócios e escolha de *portfólios* de novos serviços. Neste momento, a engenharia estará adequada às tecnologias de redes híbridas deixaremos, então, de ver a rede síncrona

tradicional tentando levar o tráfego de dados e/ou a rede assíncrona tentando criar uma rede de circuitos. Estaremos, enfim, dando um importante passo na visão da convergência das redes e dos serviços com a integração das diferentes camadas, resultando em oferta de novos serviços com qualidade de serviço trabalhada a cada aplicação, como apresentada na Figura 19.

Os serviços de telefonia fixa, ainda deverão ser dominante mesmo com a atual competição, as Empresas devem aproveitar o *marketing* das promoções para estimular novas vendas e, com uso de novas tecnologias de acesso, como as redes GPON, devem deixar de ver a telefonia IP como uma ameaça, pois esta estará incorporada à novos planos comerciais.

Os serviços de dados corporativos podem ter melhor enfoque motivado pelo crescimento do comércio eletrônico e novas aplicações multimídia, devido aos serviços de valor adicionado, como VPNs, extranets, comércio eletrônico, etc., que poderão gerar mais receitas. Para os serviços de dados residenciais ainda temos o acesso com tecnologia xDSL, levando maiores velocidades e expansão da *Internet* [43]. Ainda nesta classificação temos os serviços de dados móveis que utilizam as tecnologias 3G com maiores bandas passantes, tráfego de mensagens de texto e multimídia, acesso à *Internet*, etc.

Os serviços de vídeo, com oferta de alta definição para TV (HDTV) e interatividade, como vídeo sob demanda, distribuição de dados e conteúdo, dependerá muito do preço e da variedade e da qualidade do conteúdo oferecido. No próximo capítulo estaremos detalhando as redes híbridas e os serviços em teste em redes acadêmicas.

COMPARAÇÃO DA EFICIÊNCIA DAS TECNOLOGIAS BÁSICAS DE TRANSPORTE					
Esforço \ Tecnologia	SDH NG	ETHERNET	IP	MPLS-TP	PTN
TDM	alto	baixo	baixo	médio	baixo
OAM	alto	baixo	baixo	médio	alto
QoS	alto	baixo	médio	alto	alto
INVESTIMENTO	alto	alto	médio	médio	médio
PADRONIZAÇÃO	sim	sim	sim	sim	draft

Figura 21. Tabela comparativa

Analizando a tabela apresentada na Figura 21, nos perguntamos quando utilizaremos cada uma das tecnologias citadas, mas podemos avaliar pela ótica de confiabilidade, custos e capacidades necessárias a rede.

Então conseguimos mapear algumas aplicações onde cada tecnologia se adapte melhor. Exemplo: serviço LAN-to-LAN na tecnologia *ethernet*; conectividade de Dslam na tecnologia PTN; distribuição de OLTs na tecnologia MPLS-TP.

A tecnologia SDH nas redes de transportes tem elevado índice de OAM, provisionamento automático, sincronismo, etc. e utiliza também da tecnologia DWDM, otimizando os cabos ópticos instalados, levando ao aumento exponencial da capacidade de fornecimento de banda. A tecnologia MPLS-TP tem as características da tecnologia SDH com engenharia de tráfego e escalabilidade a custos menores.

Com o crescimento do tráfego da *Internet* e redução do custo dos equipamentos *ethernet* para as redes dos clientes, percebemos uma mudança no cenário das redes que levam a uma migração da infraestrutura das redes TDM para as redes de pacotes.

Nesta rede de próxima geração de transporte temos então as seguintes características (RFCs 4197/5085): orientado a conexão, comutação menor que 50 ms, OAM e gestão. Neste caso se enquadram as redes híbridas de pacotes que não vão depender do roteamento IP, com algumas funções do MPLS simplificadas, como a retirada da funcionalidade PHP, mescla de caminhos comutados por rótulos e multicaminho de mesmo custo.

Essas redes (PTN) usam as funcionalidades do MPLS-TP que atendem aos requisitos operacionais das redes e dá suporte as arquiteturas PWE3, MPLS, GMPLS. Por ser uma tecnologia orientada a conexão é agnóstica aos tipos de serviços, independente da camada e é uma plataforma multisserviço, pois, transporta serviços *ethernet* e TDM.

O *pseudowire ethernet* permite o transporte sobre estas redes e o seu encapsulamento é definido na RFC 4448. É um mecanismo que emula os atributos de um serviço não-IP em uma rede de comutação de pacotes.

2.10.2. NOVO MODELO DE REDES

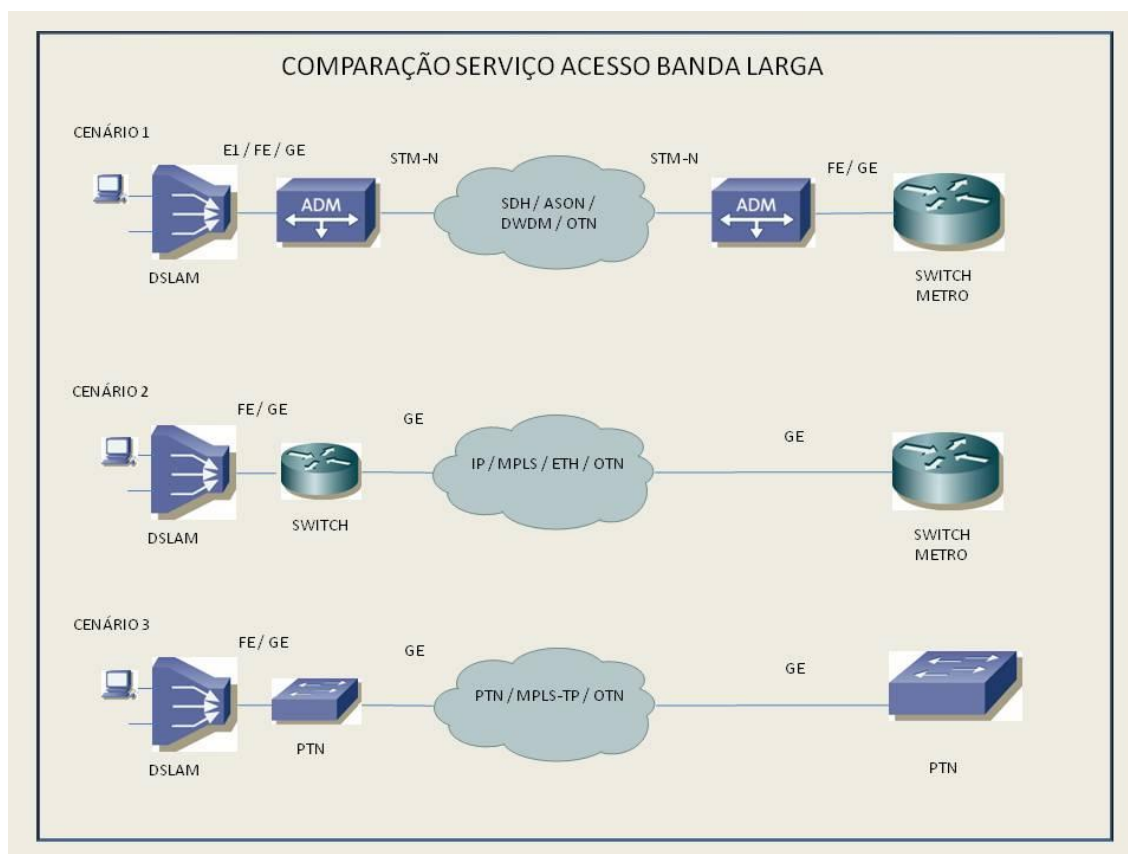


Figura 22. Comparação dos modelos de redes

A Figura 22 apresenta como podemos mapear um serviço de acesso banda larga nas redes atuais e em uma rede híbrida.

O cenário 1 apresenta uma rede de transporte tradicional, orientada a conexão. Na interligação da rede da operadora com o cliente, temos o equipamento de dados com interfaces E1, *Fast ethernet* ou *Gigabit ethernet*, interconectando com o equipamento ADM. Estas redes apresentam controles robustos de proteção e restauração, sincronismo, etc., no entanto, são ineficientes para transporte de tráfego de dados, pois mesmo com o uso de técnicas como GFP e VCAT, existe a restrição que os VCs devem ter os mesmos tamanhos. Os sistemas de provisionamento e detecção de falhas são distintos.

O cenário 2 apresenta uma arquitetura de rede não orientada a conexão, alto *overhead*, e a existência de diversos protocolos que aumentam a complexidade, tanto do provisionamento como em casos de falhas, para realizar a convergência da rede. As funcionalidades de *fast reroute* e *LSP merge* requerem plano de controle de camada 3. Além disso, as *switchs* precisam ser do tipo *carrier class*, que adicionam funcionalidades para as redes metropolitanas, elevando seus custos.

O cenário 3 apresenta uma rede híbrida, onde a oferta de interfaces para o cliente aproveita a rede legada ATM, E1, *Fast ethernet* ou *Gigabit ethernet*, de forma a tratar os tipos de tráfego em matriz agnóstica. Utiliza a tecnologia MPLS-TP, que simplifica as funcionalidades da rede MPLS convencional, reduzindo os protocolos de camada 3. Utiliza uma única gerência para provisionar os circuitos como *pseudowires*. Utiliza, também, as padronizações do MEF para oferta de serviços e plano de controle GMPLS e ASON. Possui os principais controles de proteção e alarmes, como o SDH.

Pode-se então, verificar nos estudos deste capítulo, que para a oferta de serviços *e-science*, os equipamentos híbridos possuem o melhor dos mundos estatísticos e determinísticos, agregando o conceito da tecnologia ASON GMPLS (ainda em padronização), que permite dentro das especificações de QoS, em caso de falhas, redefinir as rotas, pois a arquitetura ASON já está bem integrado as tecnologias DXC, OTN e PTN.

2.11. RESUMO DO CAPÍTULO

Apresentamos neste capítulo uma variedade de tecnologias de rede de transporte, descreve os fundamentos e as vantagens de cada uma em relação ao transporte IP e como estão evoluindo: novos tipos de equipamentos, arquiteturas, tendências, etc., além das características para uma rede de nova geração e os aspectos a serem considerados na migração das redes atuais para uma rede convergente.

Conforme descrito neste capítulo, por anos os organismos de padronização (IETF, ITU e IEEE) desenvolveram suas tecnologias e arquiteturas de rede e expandiram seu escopo de aplicação, para atender à demanda emergente de uma rede de transporte de pacotes. As tecnologias de serviços de rede de camada 3, IP e MPLS, especificadas pelo IETF, foram expandidas para serviços orientados à conexão.

A tecnologia *ethernet*, inicialmente desenvolvida para o ambiente corporativo, foi estendida pelo IEEE com camadas adicionais de OAM para melhorar a sua escalabilidade e gerência. A arquitetura, OAM e proteção especificados pelo ITU-T para as redes de camada 1, como SDH e OTN, foram aplicadas às tecnologias MPLS e *ethernet* para torná-las “*transport grade*”, ou seja, com características e desempenho de rede de transporte. Estas melhorias ainda estão em desenvolvimento [10].

A convergência combinada com a economia, traz múltiplos papéis a serem avaliados, como por exemplo: os usuários; os provedores de conteúdo e de serviço e os provedores de rede. Como consequência vemos aparecer neste cenário uma nova taxonomia de redes e de serviços, mais adequada para a presente discussão, ou seja, a questão dos serviços no ambiente da convergência passa a ser vista como uma combinação de serviços com valor agregado ao usuário, unindo em uma plataforma de rede a oferta de conversação, recuperação, distribuição, etc.

3. REDES HÍBRIDAS E SERVIÇOS

3.1. CONCEITOS

Segundo Sunan Han[14], redes híbridas são a combinação de tecnologias que tornam as redes de transporte mais econômicas. Os esforços em construir uma rede híbrida colocam como desafio para os projetistas de redes, a integração das diferentes plataformas, como a tecnologia SDH e a DWDM, que foram projetadas de forma independente. O esforço para realizar esta integração começou com o AOW (*ADMon-a-wavelength*) centralizando uma matriz STS no equipamento ROADM.

Bijan Jabbari[71] cita que as novas aplicações de redes tendem a ser construídas sobre recursos de redes híbridas. Apresenta uma implementação baseada em política de gerenciamento e aprovisionamento baseado no projeto DRAGON, que um dos objetivos é avaliar os problemas da arquitetura, como a abstração do domínio NARB (*Network Aware Resource Broken*), como criar circuitos híbridos interdomínios. Ainda menciona que uma rede híbrida é a utilização de uma mesma infraestrutura de comunicação para prestar serviços tanto de pacotes IP e de circuitos.

Uma rede híbrida permite o uso de diferentes tecnologias, em uma mesma infraestrutura. A sofisticação das aplicações multimídia distribuídas sugere alto grau de complexidade exigido dos serviços de comunicação e as dificuldades associadas aos projetos dessas aplicações. Grande parte das dificuldades consiste em oferecer serviços capazes de manter as características necessárias a diversidade de tipos de fluxo. Isto torna cada vez mais difícil encontrar uma arquitetura única, genérica o suficiente, para acomodar os diversos serviços e adaptações através de simples parametrizações ou opções pré-definidas.

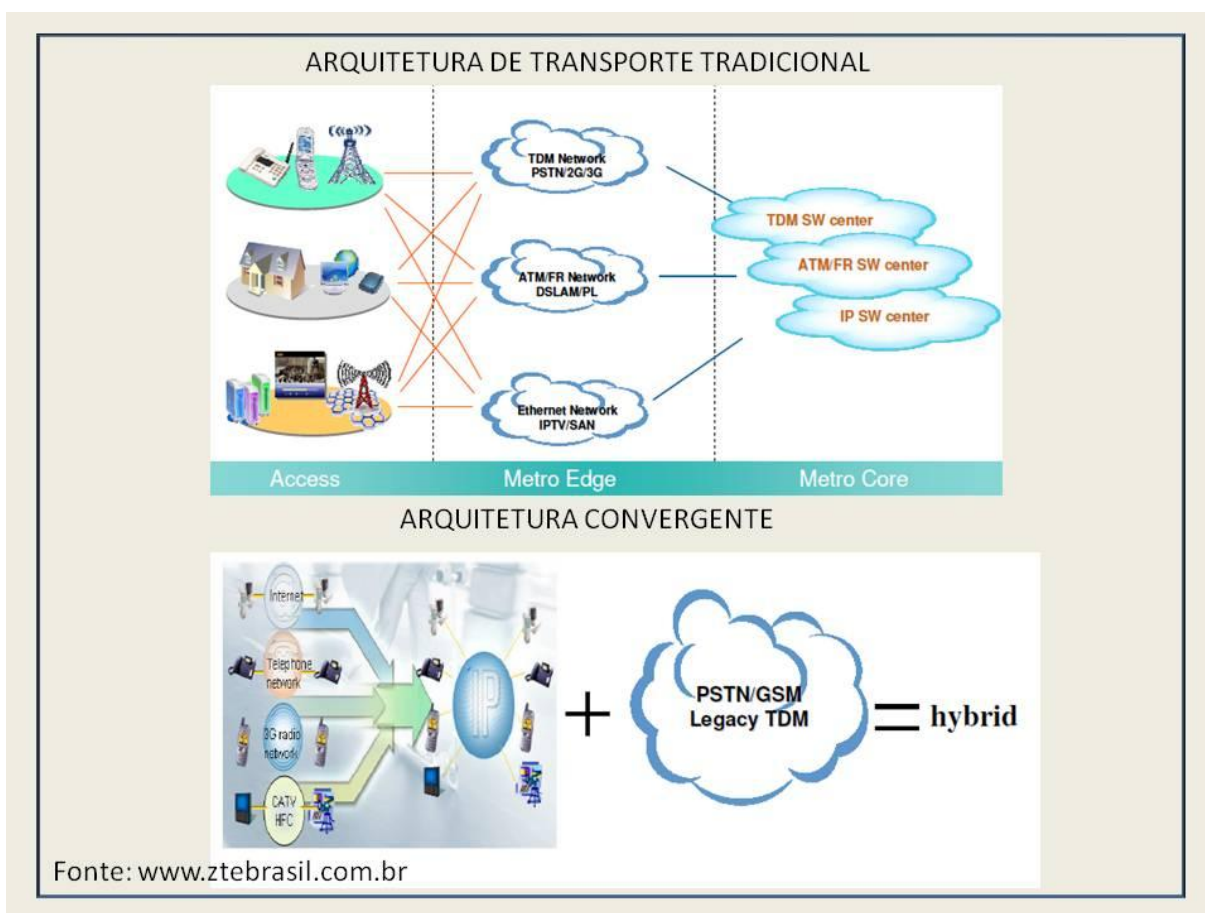


Figura 23. Nova Arquitetura

Conforme apresentado na Figura 23, uma rede híbrida permite o uso conjunto, em uma mesma infraestrutura de rede, de diferentes tecnologias. Por exemplo, é possível controlar e configurar em conjunto comutadores WDM, *switches ethernet* e roteadores IP/MPLS, através de um plano de controle. Assim, para cada perfil de cliente e aplicação, é possível alocar os recursos desejados, usando uma mesma camada de controle e os mesmos insumos operacionais. Para isto é usado o protocolo GMPLS na camada de controle, que faz o papel de integrar e unificar a operação dos vários tipos de equipamentos, com diferentes capacidades de comutação. Na sua essência, redes híbridas são redes que proveem múltiplos tipos de serviços integrados em uma mesma infraestrutura, por exemplo, uma rede que prove roteamento IP e serviços *lightpath* (banda larga) [14].

A necessidade de compartilhar informações entre pessoas tem se tornado tão importante que provoca um aumento significativo do tráfego de dados, voz, vídeo e *Internet*. Este crescimento vem criando um mercado competitivo e a necessidade de novos serviços, levando à integração das redes tradicionais. Porém, hoje estas redes ainda não são assim. Os

maiores desafios que temos que enfrentar é a convivência com novas tecnologias e atendimento aos serviços tradicionais, que ainda demandam recursos nas redes das operadoras de telecomunicações, com o a rede legada, além dos sistemas de gerência existentes ainda não serem fim a fim na criação de circuitos.

Tradicionalmente as redes de comunicação de dados podem ser divididas em redes que oferecem canais dedicados e as que utilizam técnicas de comutação de pacotes.

Esta estanqueidade vem diminuindo com o uso de tecnologias que permitem circuitos virtuais e emulação de circuitos. Entretanto, a configuração de um circuito virtual em uma rede de pacotes depende de um esforço considerável, sendo necessário o envolvimento direto da equipe de configuração. Cada tipo de equipamento tem a sua própria técnica para implementar a criação de circuitos virtuais, e na maioria das vezes, ainda com criação de circuitos estáticos. O gerenciamento em caso de falha na rede leva a necessidade de analisar o circuito em etapas (*troubleshooting*) até conseguir identificar em qual ponto da rede ocorreu o evento, isto na prática, envolve equipes também diferentes.

A crescente demanda por serviços multimídia e aplicações em tempo real torna fundamental a necessidade de estudar qualidade de serviço. As redes IP provêm entrega dos dados baseado no melhor esforço (*Best Effort*) e esta estratégia permite que a complexidade fique nas máquinas periféricas (*end-points*), permanecendo a rede relativamente simples. Isto explica o suporte ao seu crescimento, porém levando à degradação do serviço. O aumento de banda é o primeiro passo para acomodar as aplicações em tempo real, mas ainda assim, não evita o *jitter*.

Para se prover serviço adequado e algum nível qualitativo e quantitativo, deve ser acrescentado algum nível determinístico. Isso significa a adição de inteligência à rede para que possa distinguir o tráfego com requisitos em tempo real, daquele que pode tolerar atraso, *jitter* ou alguma perda. É isso que fazem os protocolos de QoS, permitindo diferenciar os diferentes fluxos na rede e reservar uma parte da banda para os que necessitam de um serviço contínuo. Existem maneiras de caracterizar o QoS, como sendo a habilidade de um elemento de rede prover algum nível de garantia para a entrega consistente de dados na rede [4].

As redes híbridas, em especial aquelas com comutação dinâmica, prometem o melhor dos dois mundos com importantes ganhos operacionais. Recentemente, algumas aplicações científicas demandam a transmissão de vídeos com alta resolução e altas capacidades, muitas vezes com aplicações de computação distribuída em *grid* ou acessos a instrumentos remotos.

E como as redes vão se comportar com estas novas demandas? Serão necessário aumento das capacidades dos enlaces, garantias de QoS, com isolamento do tráfego e caminhos dedicados, circuitos pré-estabelecidos e aprovisionamento dinâmico, via sinalização ou gerenciamento. Algumas destas questões estão representadas na Figura 24 a seguir e são objetos de estudo do capítulo 4.

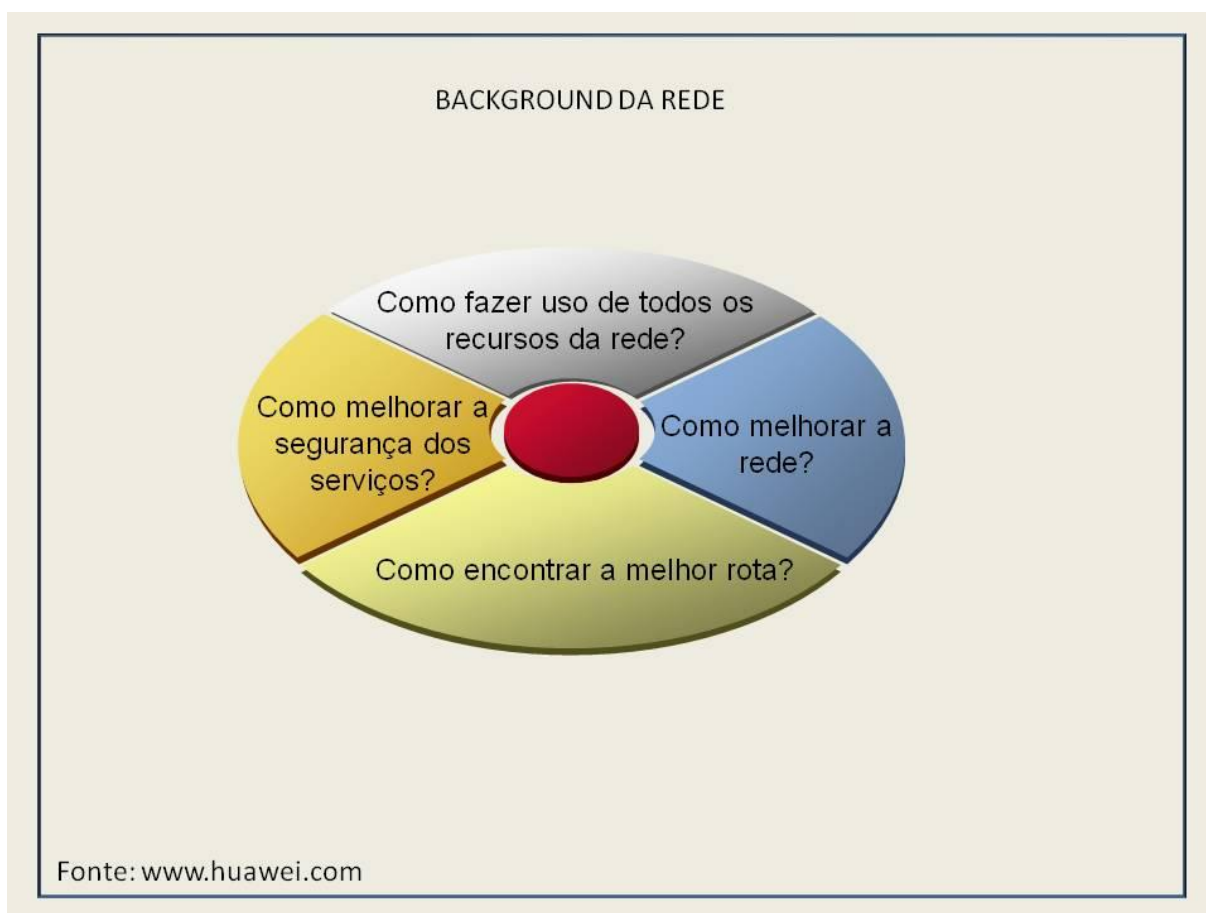


Figura 24. Background da rede

Neste capítulo serão detalhados os tipos de redes híbridas que estão sendo implantadas e suas referências de serviços testados, os conceitos de rede DCN, serviços experimentais em redes acadêmicas, além dos requisitos de QoS, TE, proteção e restauração.

3.1.1. VANTAGENS REDES HÍBRIDAS

Com o rápido crescimento dos serviços de pacote, a rede de transporte SDH que se baseia no canal TDM de granularidade rígida e fixa, encontra desafios em algumas áreas, tais como: o serviço de aprovisionamento e processamento, largura de banda, eficiência e

escalabilidade da rede. Para resolver estes problemas, uma plataforma de transporte de pacotes, ou seja, híbrida, que visa à camada de agregação da rede de transporte móvel para a evolução *all-IP* e rede LTE (*Long Term Evolution* - 4^o G) torna-se interessante.

A arquitetura híbrida é aplicável ao transporte do futuro *all-IP* e atendendo às exigências de novos serviços. Vai utilizar as seguintes tecnologias: PWE3 (*pseudowire* - RFC 3916 e 3985) que é orientado para acesso ao serviço e assegura a qualidade de transporte semelhante aos tradicionais serviços TDM / ATM / *Ethernet*, *Carrier-class* com mecanismo de sincronização do relógio SDH fim a fim na prestação de serviços, OAM (RFC 6371) e serviço de proteção, QoS hierárquico e orientado a conexão.

A tecnologia aumenta a flexibilidade da rede e a eficiência do transporte, e estruturalmente reduz o custo total de investimento. Suporta o tradicional TDM com sincronização do relógio, e outras soluções de múltiplos *clocks* como síncrono *ethernet*, recuperação de relógio adaptativo (ACR), e IEEE 1588v2. Estas soluções de relógio para atender aos requisitos de sincronização de rede linear, anel e de rede *mesh*. Interfaces de serviços tradicionais, tais como E1 (TDM/IMA), IP sobre E1, TDM STM-1, ATM STM-1, POS STM-1/STM-4 interfaces *ethernet*, tais como, FE, GE e 10GE. Já QoS pode ser analisado pelo lado do usuário e da rede. Assim, temos QoS-TE no lado da rede e as políticas de QoS hierárquicas no lado do usuário para alocar e supervisionar largura de banda por fluxo de serviço, o usuário, ou par de usuários, e por porta física.

3.2. REDES HÍBRIDAS ACADÊMICAS

No Brasil foi criada a rede híbrida chamada Cipó com uso da ferramenta OSCARS, uma rede *overlay* à rede Giga, com *switches* de camada 2 oferecendo conexões ponto a ponto. Em 2011 a RNP iniciou a operação do Serviço Experimental de Circuitos Aprovisionados Dinamicamente (SE- Cipó). O objetivo é possibilitar que circuitos fim a fim (ou *lightpaths*), sejam provisionados na rede Ipê sob demanda de maneira automática, em substituição ao tradicional provisionamento manual, uma atividade lenta, trabalhosa e sujeita a erros. A ideia é que esses circuitos possam ser facilmente agendados pelos usuários por meio de uma interface *web* [68].

Entre as aplicações beneficiadas pelo provisionamento dinâmico de circuitos estão aquelas que requerem qualidade de serviço na forma de prioridade de tráfego de dados e as que precisam de grande largura de banda, como ocorre em aplicações “intensivas em dados” das áreas de meteorologia e física de altas energias experimental, por exemplo [5].

A RNP estudou alternativas para a arquitetura e forma de implantar o serviço no *backbone* Ipê e / ou nos clientes. Fez a opção por implementar apenas no seu *backbone*, segregando o tráfego e utilizando protocolo MPLS e nos clientes utilizando componentes externos (VLSR - *Virtual Label Switch Router* e NARB) que realizam as configurações de controle do equipamento, ambos os ambientes com o componente que realiza as configurações interdomínio (OSCARS) e para solicitar os circuitos, o usuário deve usar a interface MEICAN (*Management Environment of Inter-domain Circuits for Advanced Networks*), que é um ambiente de gerenciamento de rede DCN.

MEICAN é um sistema *web* para gerenciamento de reservas com interface amigável, desenvolvido em PHP pelo grupo de redes da UFRGS [68], que opera interligado ao OSCARS através de API. Os usuários acessam o MEICAN, via o site da RNP, informam o ponto de origem e de destino, banda e o tempo que o circuito deve se manter ativo, e que vai interagir com o OSCARS, que receberá a solicitação do usuário e acessará os roteadores do *backbone* da RNP.

A expectativa é que o serviço de provisionamento dinâmico de circuito desenvolvido pelo SE-Cipó (nome criado para o serviço experimental) esteja disponível para todos os usuários da RNP em 2013. Ainda estão em discussão as regras para adição de políticas, escolha do caminho do circuito através de mapas e integração com ferramenta de monitoração, bem como as aplicações que poderão ser oferecidas aos clientes e também o teste com a interface NSI (*Network Service Interface*) que vai fornecer uma interface genérica de rede de serviços como usuários finais, para definir a troca de mensagens entre cliente / servidor, hoje oferecido pelo IDC (*Inter-Domain Controller*).

A partir disto, a RNP terá uma arquitetura híbrida, onde uma mesma infraestrutura de comunicação é usada simultaneamente para prestar serviço de roteamento de pacotes IP e, adicionalmente, o provisionamento de circuitos fim a fim para aplicações selecionadas. O SE-Cipó prevê interoperabilidade com os serviços correspondentes de parceiros internacionais [68].

Alguns programas de pesquisas nos Estados Unidos, Europa e Ásia, nas redes acadêmicas de diversas universidades, vêm buscando uma redução na complexidade do gerenciamento e aumento da confiabilidade das atuais redes, usando o protocolo GMPLS como nova forma de prover serviços. O plano de controle GMPLS pode ser integrado com aplicações para realizar o provisionamento dinâmico de circuitos sob demanda, implementando o conceito de redes dinâmicas. A DCN é uma arquitetura de rede que permite o desenvolvimento de serviços para criação de circuitos ponto a ponto entre usuários finais com banda dedicada para algumas aplicações, usada em projetos como LHC (acelerador de partículas [70]) e em outros projetos *e-science* com plano de controle GMPLS.

Em 2009 a rede híbrida DCN Internet2 (EUA) [48] passou a oferecer o serviço DCN/ION (DCN *On-Demand*) que utiliza um plano de controle GMPLS, separado do plano de dados, para instanciar circuitos. A Internet2 usa a versão *open source* da DCN *Software Suite* (DCNSS), que contém os aplicativos DRAGON (*Dynamic Resource Allocation via GMPLS Optical Networks*) e OSCARS (*On-demand Secure Circuits and Advance Reservation System*), para implementar as funções do plano de controle, com equipamentos DWDM da Ciena [71].

Também a rede híbrida DCN Geant (Europa) desenvolveu uma solução utilizando o GMPLS como plano de controle, o AutoBahn, realizando testes com equipamentos Alcatel 1678 (ASON) em alguns POPs (*Point of Presence*) em comprimentos de onda dedicados). A principal diferença entre o AutoBahn (*Automated Bandwidth Allocation across Heterogeneous Networks*) e o DCNSS é o controle centralizado dos elementos da rede, abordagem adotada pela Geant [45].

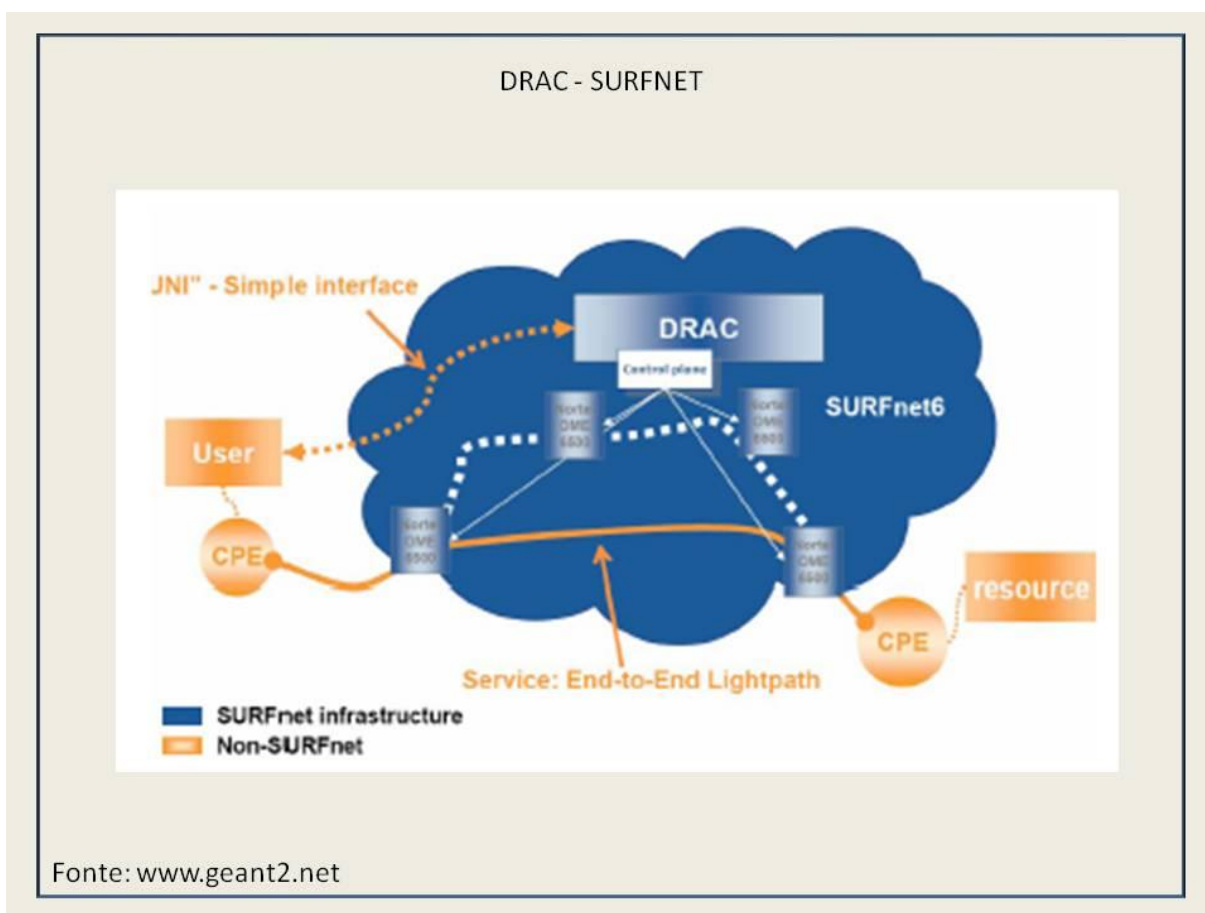


Figura 25. Projeto DRAC

Algumas outras redes que também buscam oferecer serviços de redes híbridas são a Surfnet (Holanda), Canarie (Canadá) e a Esnet (EUA) [45], que oferecem serviços de circuitos fim a fim em camada 1 e 2. A Surfnet é baseada em fibra apagada e utiliza um sistema de gerenciamento próprio. Utiliza equipamentos Nortel Network Corporation que desenvolveu o plano de controle DRAC (*Dynamic Resource Allocation Controller*), que aproveita o plano da rede existente. Todas as suas fibras ópticas são iluminadas com a tecnologia DWDM, como apresentado na Figura 25 [45].

A rede híbrida Canarie desenvolveu a solução UCLP/Argia (*Controlled LightPaths*), com uma abordagem diferente das demais ao usar o plano de gerência para fazer o provisionamento dinâmico dos circuitos, com equipamentos ROADM. O UCLPv2 é um *software* que dispõe para os usuários criar suas próprias topologias. Utiliza o arcabouço IaaS (*Infrastructure as a Service*) no controle dos recursos de rede. A Esnet, similar à Internet2, utiliza o *software* OSCARS [45].

ALGUMAS NRENS			
LOCAL	REDE	PLANO CONTROLE	VENDOR
BRASIL	REDEH	CIPO - OSCAR	JUNIPER
CANADA	CANARIE	UCLP - CANET4	CISCO E NORTEL
EUA	INTERNET2	DRAGON - GPENI	CIENA
EUA	ESNET	OSCAR - HOPI	NORTEL
EUROPA	GEANT2	JRA3 - AUTOBAHN	ALCATEL
HOLANDA	SURFNET6	DRAC	NORTEL

Figura 26. Redes Híbridas

A Figura 26 traz um resumo de algumas redes híbridas já implementadas. Das diversas redes híbridas, o plano de controle intradomínio são os G-LAMBDA no Japão e OSCARS nos EUA, os demais todos são interdomínio. O G-LAMBDA é um projeto de pesquisa para estabelecer interface padrão entre o *grid* e a rede [45]. Estas iniciativas estudam soluções que viabilizem a implantação de redes e serviços, com abordagens de planos de controle / gerenciamento, ambos customizados e testar interfaces com demais planos de controles, fornecendo serviço confiável de dados em alta velocidade e com criação de *middlewares* que tornem possível aos clientes aprovisionar o serviço dinamicamente [45].

A criação dinâmica de circuitos fim a fim exige que as redes sejam capazes de fazê-lo, uma das soluções seria deixar a criação dinâmica dos circuitos nas camadas N1 e N2 apenas no núcleo e deixar o roteamento N3 para as bordas. Neste caso os roteadores de borda ficam responsáveis pelo aprovisionamento dinâmico de circuitos.

-Vantagem:

- Como requisitos as redes oferecem banda larga, baixo atraso e facilidade de alocação de recursos de rede.

-Desvantagem:

- Alto custo de roteadores camada 3.

-Necessidade de mecanismos eficazes de alocação de recursos computacionais.

Na 11ª edição do *Global LambdaGrid Workshop* (GLIF2011), realizado no Rio de Janeiro em setembro de 2011, foi demonstrada a interoperação entre o SE-Cipó e o ION, serviço semelhante da rede norte-americana Internet2, no contexto do projeto DYNES (projeto colaborativo da NSF sobre microestruturas da decomposição de revestimentos produzidos pelo plasma), dentro do programa *Internet Research Networks Connections* (IRNC [69]) da *National Science Foundation* dos EUA, a partir de 2010 [68].

Na sexta conferência internacional de IP + Rede Óptica (iPOP 2010) realizada em junho de 2010 em Tóquio, foram apresentadas novas aplicações, serviços e diversos estudos que vem sendo realizados para verificar um caminho que atenda a evolução com menores custos para os provedores [49].

O jornal ONTC (IEEE *ComSoc Optical Networking Technical Committee*) publicou em março de 2010 o potencial para o protocolo GMPLS para futuras aplicações [53].

O GMPLS é um conjunto de protocolos que permite descobrir recursos, aprovisionar serviços e recuperar, em caso de falhas, a rede de forma automatizada. Impulsionado pelo benefício de melhoria da confiabilidade da rede (através da proteção e recuperação de falhas) e com redução do OPEX (*Operational Expenditure*) da rede, está sendo implantado nas redes metropolitanas e em *backbones*. Em curto espaço de tempo essas redes vão permitir as operadoras de rede, fornecer serviços de valor agregado como, por exemplo, o BoD (*Bandwidth on Demand*), e as redes GMPLS terão potencial para transportar grande variedade de serviços.

Desde 2001, vários programas de pesquisas, buscando uma redução na complexidade do gerenciamento e aumento da confiabilidade das suas redes, vêm usando no protocolo GMPLS uma nova forma de prover serviços. O plano de controle GMPLS é geralmente integrado com a aplicação para realizar circuitos sob demanda, de modo que as aplicações de dados podem ser vistas como tubos com largura de banda dedicada de forma mais eficiente.

Em estudos recentes, Vicent Chan[7] argumenta que o GMPLS pode ser uma das tecnologias que permita realizar grande escala, serviços dinâmicos em comprimento de onda (*lambda*) para massas. Ao combinar a multiplexação estatística de acesso de redes e o aprovisionamento de *lambda* no *core*, os pesquisadores tentam tornar a *Internet* mais rápida e mais barata.

Chan [7]apresenta o desenvolvimento de protocolo de gerenciamento de rede, a utilização da comutação óptica e propõe uma arquitetura de transporte óptico (OFS) que irá explorar a comutação óptica, roteamento e tecnologias de transporte que continuem a reduzir os custos, permitindo acesso aos serviços de alta taxa para massa, mais rápido que as tendências atuais.

Na arquitetura OFS os clientes solicitam *lighpaths* fim a fim por tempo superior a 100 ms e se for concedido, os recursos dedicados serão liberados dos outros usuários que tiveram suas transferências concluídas. Uma vez que não poderão ser atendidas todas as solicitações, haverá uma compensação entre os indicadores de desempenho, como atraso, probabilidade de bloqueio e utilização do comprimento de onda. Este *trade-off* (conflito de escolha) empresta qualidade de serviço que existe na rede e isto pode ser implementado incorporando, funções de custo adequado no protocolo de agendamento.

Para programar as transmissões através da WAN (*Wide Area Network*), os usuários se comunicam via plano de controle com os processadores de agendamento das suas respectivas MANs (*Metropolitan Area Network*). Esses processadores coordenam a transmissão de dados pela WAN pelo plano de controle no OFS. No caso em que muitos usuários possuam agendamentos que não sejam suficientemente grandes para justificar seu próprio *lambda*, eles podem ser multiplexados para transmitir como formação de grupo *broadcast* dinâmica [7].

Spadaro [55] apresenta o projeto *Phosphorus* [13]que define a construção de plano de controle de rede (NCP - *Network Control Plane*) que integra um mecanismo de aprovisionamento GNS (*Grid Network Service*) [55]e permite a alocação e reservas de recursos computacionais de rede. O objetivo é facilitar a dinâmica do aprovisionamento automático e que inclui protocolos e mecanismos de divulgação dos recursos [57].

A convergência introduz múltiplos atores neste cenário de prestação de serviços, e reduz a margem de lucro no segmento de transporte. Isto dificulta a adoção de modelos de negócios tradicionais. Muitos dos novos serviços oferecidos mudam o esquema de comercialização e cobrança e até mesmo a garantia do serviço.

A ênfase deverá ser no próprio mercado, corporativo, massa ou de pesquisas, que naturalmente vão demandar soluções a preços reduzidos [43]. Através de parcerias, as operadoras deverão buscar atuar, em nova cadeia de valor. Novos esquemas de cobrança através de tarifas fixas (*flat rate*) costumam estimular o consumo dos serviços.

A topologia em malha é o meio mais eficiente de prover conectividade entre as redes. Os equipamentos OXCs (*Optical Cross-Connect*) proveem os canais ópticos com as mesmas taxas dos roteadores e que pode ser conectada a sistemas DWDM, esta geração de equipamentos prove compatibilidade com o legado existente [17].

Todas estas iniciativas visam oferecer serviços dinâmicos em redes híbridas e tudo isso vem se tornando possível devido ao desenvolvimento de uma versão da interface do cliente UNI 1.0, existente hoje nos equipamentos ópticos, em diversas pesquisas em redes acadêmicas com a participação de fabricantes, com a colaboração do OIF, para a evolução da nova interface de sinalização UNI 2.0, que estaremos descrevendo melhor no capítulo 4 [22].

3.2.1. REDES DCN

Segundo detalhado por Jabbari [15] o *software* DRAGON foi desenvolvido por um projeto com o mesmo nome financiado pelo *National Science Foundation* (NSF), que tem a proposta inicial de desenvolver e integrar as tecnologias (DWDM, SDH, MPLS), incluindo *hardware* e *software*, para prover o serviço dinâmico de *lightpaths*. O *software* desenvolvido é parte da DCN *Software Suite* e é utilizado como *Domain Controller* (DC) pela Internet2 no serviço ION (*Interoperable On-Demand Network*). O DRAGON permite o controle dinâmico de múltiplas tecnologias no plano de dados usando o protocolo GMPLS, bem como adaptações do RSVP-TE e OSPF-TE para a reserva e computação de circuitos na rede. Seus principais componentes são:

- VLSR – responsável por controlar os elementos de rede para que estes sejam capazes de interagir via GMPLS.

- NARB / RCE – inclui a entidade responsável pela computação dos caminhos na rede pelos quais os circuitos virtuais irão passar o PCE (*Path Computation Element*).

O OSCARS *Inter-Domain Controller* (IDC), desenvolvido pela ESnet, com financiamento do DOE *Office of Science*, também é parte da DCN *Software Suite* utilizada

pela Internet2. O OSCARS, na forma como foi implementado na Internet2, é responsável pela sinalização e computação de caminhos interdomínio. Todavia, se os elementos de rede forem capazes de utilizar um plano de controle GMPLS, pode ser usado sem a integração como o DRAGON. Neste caso o OSCARS fará a computação também de caminhos intradomínio [54].

O AutoBahn faz parte do projeto de evolução da rede Geant, utilizando tecnologias de transporte para oferecer serviços de estabelecimento de circuitos dinâmicos, além de serviços baseados em IP. O AutoBahn é uma plataforma que fornece uma interface para a criação de circuitos dinâmicos virtuais via plano de controle. Seus componentes principais componentes são:

- cNIS – responsável por prover um repositório unificado com informações sobre o domínio administrativo e pela análise da topologia de dados.

- Topology Proxy* (TP) – faz a interface com os elementos de rede, de forma semelhante ao VLSR.

- User Client* – faz a interface com o usuário.

Ambas as plataformas estão ainda em fase de testes pela RNP, entretanto os serviços aqui descrito, podem ser implementados com qualquer uma das soluções, desde que atendam à funcionalidade básica de agendamento, estabelecimento e remoção de circuitos dinamicamente. As demais facilidades ofertadas, em especial as questões envolvendo QoS e TE podem ser implementadas usando, por exemplo, a infraestrutura de rede do fornecedor Juniper prevista para o novo *backbone* da RNP. Aliás, a capacidade de realmente configurar os requisitos de QoS e TE solicitados em uma reserva, depende fortemente das possibilidades dos elementos de rede usados na infraestrutura da rede do provedor.

Para monitorar o cumprimento dos requisitos de QoS oferecidos nas diversas modalidades do serviço, bem como permitir ao usuário final acompanhar dados de desempenho da rede e do serviço, deve ser usada uma ferramenta capaz de monitorar fim a fim caminhos que cruzam vários domínios. O PerfSONAR (*Performance focused Service Oriented Network monitoring Architecture*) é um sistema baseado em *web* para a coleta e publicação de dados de monitoração de desempenho de rede, desenvolvido especificamente para monitorar e resolver problemas de desempenho fim a fim em múltiplos domínios. Além de atender à demanda de monitoração deixada pelo plano de controle, o PerfSONAR também pode ser integrado ao OSCARS fornecendo serviços de topologia e localização do IDC [45].

A utilização conjunta de uma solução de DCN e do PerfSONAR sobre o *backbone* MPLS, resultará em uma plataforma capaz de prestar o serviço de comutação dinâmica de circuitos em uma rede híbrida, com qualidade de serviço e monitoração fim a fim.

3.2.2. PROJETO DRAGON

Na rede híbrida Internet2, o plano de controle é o DRAGON desenvolvido para permitir a criação de circuitos *ethernet* fim a fim, permitindo que seu acionamento seja dinâmico a partir do cliente. Este *software* foi utilizado nos testes da rede Cipó, no Brasil.

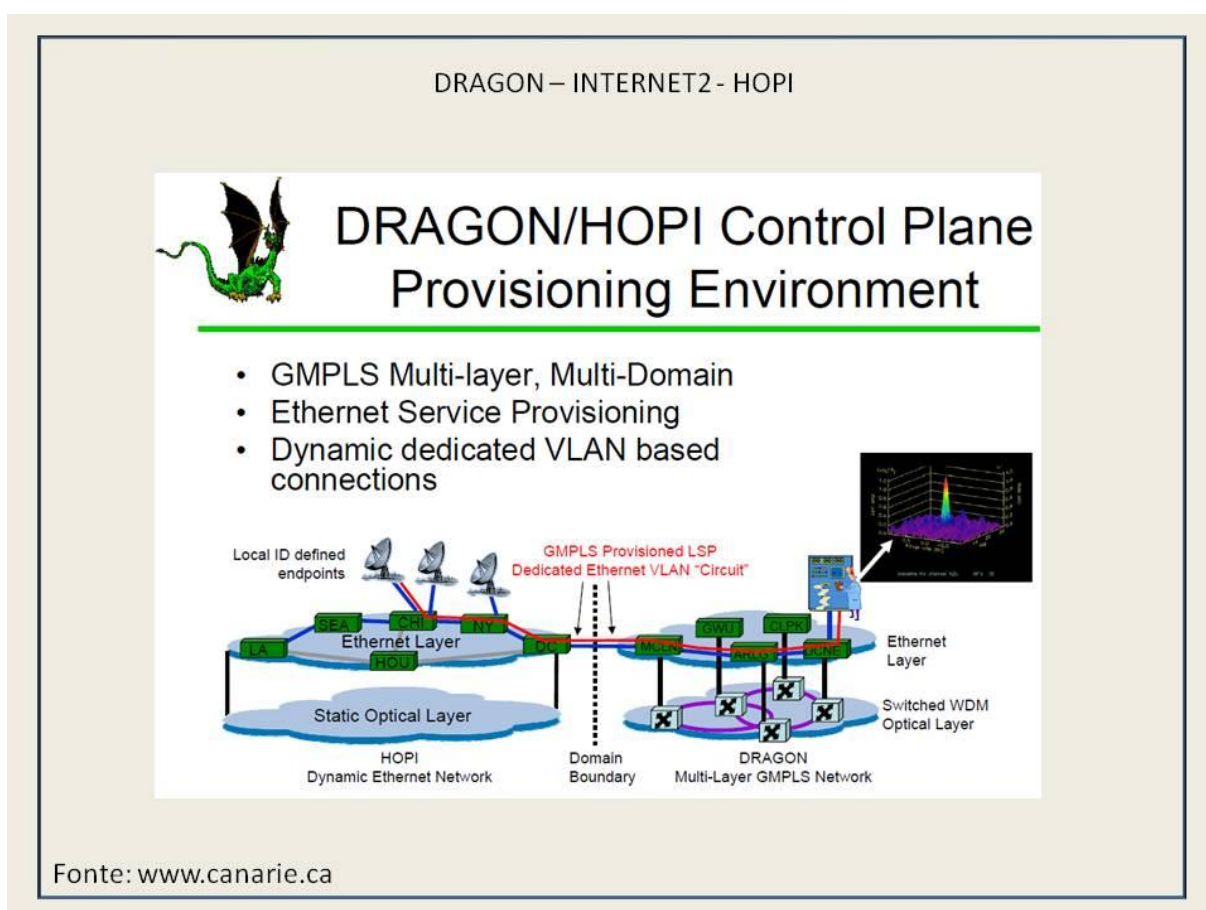


Figura 27. Projeto Dragon

O objetivo destes estudos em redes acadêmicas (Figura 27) é o atendimento a demandas de serviços, como interligar *clusters* computacionais, *arrays* de armazenagem, sensores remotos, e outros instrumentos distribuídos em escala mundial e de aplicação

específica [44]. Estes serviços de rede avançados são motivados para melhoria dos serviços como radiotelescópios, poderosos aglomerados computacionais, e grandes repositórios de dados podem ser facilmente compartilhados entre os pesquisadores, independentemente da sua localização. Outras comunidades e aplicações cujo interesse neste tipo de serviços é a missão empresarial/serviços críticos, e de construção (ou de engenharia de tráfego) de melhor esforço da rede IP.

Estes super usuários exigem o que têm sido referidos como serviços de rede "deterministas". Neste contexto, "determinísticos" implica definição de nível de serviço garantido. Estes parâmetros incluem em nível de serviço uma banda mínima. Uma habilidade para especificar latência, *jitter*, perda de pacotes, e outros parâmetros são verificados também. Além disso, estes serviços devem ser instalados em uma base interdomínio, que incluem funcionalidades de autenticação, autorização, contabilidade (AAA) e *scheduling*.

Para atingir estes objetivos o projeto DRAGON desenvolveu um código fonte aberto com o protocolo GMPLS, incluindo um plano de controle baseado em avançados serviços interdomínio, técnicas de encaminhamento e regras de formalização. Uma rede de referência foi construída na área de Washington DC (Internet2), conforme descrito por Bijan Jabbari [15].

DRAGON é um projeto colaborativo entre a Universidade de *Maryland* (UMD) *Mid-Atlantic Crossroads* (MAX), *University of Southern California Information Sciences Institute Oriente* (USC / ISI), e *George Mason University* (GMU). Utiliza o protocolo GMPLS como o alicerce para a rede básica de controle e aprovisionamento. Existem vários elementos funcionais identificados na arquitetura do plano de controle DRAGON: [43]

- Network Aware Resource Broker* (NARB).

- Virtual Label Switch Router* (VLSR).

- Client System Agent* (CSA).

- Application Specific Topology Builder* (ASTB).

O NARB é uma entidade que representa um sistema autônomo (AS). O caminho NARB serve como mecanismo de cálculo que os sistemas ou outros dispositivos poder consultar para saber sobre a disponibilidade da rede. Possui um subcomponente autônomo (*stand-alone*) chamado a *Resource Computation Element* (RCE), que realiza a tarefa de computação. O NARB também é responsável pelo roteamento interdomínio. NARBs pares

em vários domínios realizam a troca de informações da topologia que permite ao algoritmo interdomínio (*path computation*) o provisionamento do *Label Switched Path* (LSP).

Passaremos a descrever as características da ferramenta DRAGON conforme relatório de pesquisa elaborado por Ferreira A. E. e Aguiar E. C. V. P. [10].

A topologia interdomínio pode ser baseada na topologia real conforme descoberto com o protocolo OSPF-TE, ou opcionalmente com base em um ponto de vista do domínio (arquivo gerado pela configuração automática ou síntese do estado de dados OSPF *link*). Este domínio prevê mecanismos administrativos para um domínio anunciar ao mundo exterior uma visão mais simplificada da sua topologia. Isto permite que os seus domínios possam ocultar partes da topologia real, bem como minimizar a quantidade de atualizações externas exigidas. O *trade-off* é o reduzido grau de precisão para os cálculos de caminho. Cada domínio administrativo pode utilizar a configuração dos parâmetros para personalizar o seu domínio de abstração ao nível desejado.

Outro elemento do plano de controle chama-se VLSR que fornece um mecanismo para integrar os equipamentos de rede que não tenham o protocolo GMPLS. O VLSR traduz os protocolos padrão GMPLS em protocolos específicos, para permitir a reconfiguração dinâmica de dispositivos não GMPLS. O PC VLSR consiste de um plano de controle que inclui a pilha de protocolos OSPF-TE e RSVP-TE, e atua como um *proxy* para o equipamento não GMPLS. Isto permite que dispositivos não GMPLS possam ser incluídos no caminho fim a fim. O principal uso para VLSR sobre o projeto DRAGON é controlar as *switches ethernet*, através do plano de controle GMPLS. No entanto, o VLSR também foi adaptado para controlar comutadores ópticos e *switches*. Em cada *switch ethernet*, o VLSR traduz RSVP-TE que sinaliza mensagens em locais alternando comandos e criando associações VLAN-portas, juntamente com a banda garantida solicitada [30].

Sempre que um circuito *ethernet* (ou LSP) é criado ou cancelado, a informação de largura de banda e VLAN *tag* é atualizada através da distribuição de mensagens de anúncio dos *links* OSPF-TE, a fim de manter o estado da ligação adequado em toda a rede. O VLSR PC utiliza uma combinação da RFC 2674 SNMP (*Simple Network Management Protocol*) e *Command Line Interface* (CLI) de comandos para o controle local do *switch ethernet*.

O CSA é um *software* que é executado em qualquer sistema que termina o plano de dados para provisionamento de *link* de serviço. Este *software* participa nos protocolos GMPLS para permitir pedido de provisionamento do sistema cliente da ponta A para a ponta

B do cliente. Neste contexto, *Client System* (CS) é um termo muito amplo. Isto poderia incluir uma máquina, um *cluster* computacional, um roteador, um rádio telescópio, ou vários outros equipamentos de redes. Neste contexto, um “cliente” é qualquer sistema que esteja solicitando serviços à rede [71].

A arquitetura DRAGON inclui a noção de criação de topologias para aplicações específicas (ASTB). Estes são solicitados por um usuário final e é geralmente um conjunto de LSP's de um pedido para ser criado como um grupo. O elemento DRAGON conhecido como o aplicativo específico *Topology Builder* (ASTB) é o responsável pela coordenação, em resposta a este pedido. O ASTB utiliza as capacidades do plano de controle DRAGON dos elementos (NARB, VLSR, CSA) para solicitar criação dos LSP's, manter o mapeamento de cada LSP, agrupamentos específicos para topologias, e interagir com os usuários [43].

3.2.3. REDE HÍBRIDA ASON

Conforme já citado anteriormente, a configuração de um serviço é complexa e a capacidade de expansão leva um longo período. A utilização da largura de banda é de baixa eficiência e as redes com topologia anel, metade da banda será reservada para proteção. Os esquemas de proteção e de recuperação são poucos, que são detalhados no item 3.6.

Com as redes híbridas ASON a configuração do serviço será fim a fim e a rede pode realizar estas operações de forma automática e com a função de roteamento, pode prover proteção e reservar recursos, sendo que a principal topologia é a malha. A estrutura da rede será composta pelos NEs, *links* TE e domínios SPC.

Para o estabelecimento dos circuitos, o cliente através da sua interface UNI-C, demanda o início da conexão via sinalização para o elemento NCC, que vai estabelecer o circuito na rede. Poderão ser criados dois tipos de serviços, o EPL (*Ethernet Private Line*) e o EVPL (*Ethernet Virtual Private Line*), sendo o primeiro ponto a ponto e o segundo ponto multiponto [22].

O serviço pode ser criado entre diferentes redes, como SDH e a G.805 (arquitetura), que define as funções de adaptação e terminação entre o cliente e a rede. A interface UNI-C é conectada a interface UNI-N em padrão *ethernet*, no escopo do serviço EPL e EVPL [51].

Para cada serviço *ethernet* requisitado entre duas interfaces UNI de cliente, uma chamada é criada, que resulta na criação da conexão, que deve ser estabelecida após a camada de serviço for criada. Para o serviço EPL, é criado um EVC na rede entre as interfaces UNI-N e para o serviço EVPL, são criados vários EVCs, seria um serviço ponto-multiponto. Para a proposta do novo serviço, utilizamos a configuração ponto a ponto, conforme iremos detalhar a seguir.

3.3. SERVIÇOS EXPERIMENTAIS

Dentro dos cenários dessas redes híbridas, podemos destacar alguns serviços que estão sendo oferecidos nas redes acadêmicas da Internet2 e da Geant [56].

Na Europa está sendo desenvolvido o AutoBahn para estabelecer os circuitos nas bordas de uma rede híbrida, com reservas antecipadas. Para isso o sistema deverá lidar com a reserva dos circuitos através de vários domínios. O processamento interdomínio e intradomínio serão atribuídos aos protocolos IDM (*Inter-Domain Manager*) e DM (*Domain Manager*). O AutoBahn é um sistema desenvolvido pelo GN2 (*Geant Network*) que permite a criação automática e de reservas de circuitos BoD em rede híbridas, sobre múltiplos domínios e com diferentes tecnologias, com banda dedicada em camada 2, onde o usuário será capaz de emitir um pedido específico e a rede irá fornecer os recursos sem a intervenção humana.

O foco inicial do serviço era oferecer circuito *ethernet* ponto a ponto, que poderia ser nativo ou encapsulado em outra tecnologia como o SDH, o usuário final poderia ter um acesso *ethernet* de 1G ou de 10G e o serviço seria transparente. No futuro ela irá utilizar uma arquitetura de rede com NG-OTN como suporte aos serviços TDM, pacote e IP. No final do projeto GN2 o AutoBahn estava apoiado em três tecnologias, o *ethernet*, SDH e MPLS, implantado na rede da Geant2 usando equipamentos da Alcatel.

Outro exemplo é o DRAC desenvolvido pela Nortel Network Corporation na rede híbrida Surfnet que provisiona circuitos de camada 1 e camada 2, sendo um circuito entre duas portas *ethernet* atravessando uma rede SDH.

O projeto OSCARS consiste em um *front end* e *back end* desenvolvido para a rede da Esnet nos EUA, com foco na solução BoD. Inicialmente foi usado em LSPs para túneis IP e depois com VPN L2, com solução com um único domínio. O *front end* provê uma interface de usuário via *web* com autenticação e permissão, o *back end* provê a necessidade de banda.

O OSCARS é um protótipo de serviço com provisionamento dinâmico, com circuito de banda garantida e foi financiado pelo DOE (departamento de energia), é um *software* baseado no *software* BRUW (*Bandwidth Reservation for User Work* [45]) desenvolvido no projeto da Internet2, e hoje eles possuem o mesmo código base comum [54].

A Internet2 ION é um serviço que prove aos usuários habilitar circuitos através de vários domínios. Utilizando uma simples e segura interface, os clientes podem criar circuitos ponto a ponto ou ainda modificar seus pedidos de conexões, que são encaminhados a um NOC (*Network Operation Center*), que oferece um range de 55M a 10G, dependendo da aplicação.

A DCN na Internet2 permite aos usuários criar circuitos em toda a infraestrutura da rede usando o plano de controle automático. Este plano inclui o DRAGON e o OSCARS. Um membro da Internet2 se conecta a DCN através de uma interface, esta poderá criar circuito na própria rede nos EUA, ou na Esnet ou Geant. A abordagem da primeira etapa é feita de forma estática, entre o usuário e a DCN. A segunda etapa é conectar no próprio domínio DCN e então dinamicamente interligar diferentes usuários, conforme Figura 28.

COMPARAÇÃO DOS SERVIÇOS				
PROJETO	NREN	TECNOLOGIAS	TIPO REDE	INTERFACE
DRAC	SURFNET	L 1 E L2 (FUTURO)	SDH / DWDM	GUI / API
DRAGON	INTERNET2 / HOPI	ETHERNET	GMPLS	GUI / API
OSCARS	ESNET	MPLS VPN	INTRA-DOMINIO	GUI
UCLP	CANARIE	L2	INTRA-DOMINIO	GUI
AUTOBAHN	GEANT	L2	GMPLS	GUI
CIPÓ	REDE H	L2	ETHERNET	API

Figura 28. Serviços híbridos

3.4. QUALIDADE DE SERVIÇO

Uma consequência da evolução tecnológica descrita é que o conteúdo (informação) e o meio de transporte (rede) puderam ser dissociados, em outras palavras, o serviço e a rede não se confundem mais, como por exemplo, conversação de voz em PC, ler textos e *e-mails* em celulares, oferecer telefonia em rede de TV a cabo, etc. Assim podemos dizer que convergência pode ser entendida como a “capacidade de diferentes plataformas de rede servirem de veículos a serviços semelhantes”, ou seja, os serviços não estão mais ligados ao tipo de rede.

Dois importantes tópicos precisam ser mencionados, para caracterizar um processo de convergência, a qualidade de serviço e o legado tecnológico. Esses limites são dados pela natureza do serviço e como é percebido pelo usuário. Para que este possa desfrutar satisfatoriamente de um serviço, a infraestrutura de rede precisa atender a requisitos técnicos diferentes por tipo de serviço. A motivação de adotar redes com classes de serviços reside no desejo de ter uma base única para o transporte das informações, com diferentes requisitos em diferentes classes ou categorias. Tendo em mente que a qualidade de serviço é definida a partir da perspectiva do usuário e então mapeada para parâmetros técnicos. Determinadas limitações no desempenho da rede podem vir a prejudicar a aceitação do serviço pelo público.

Outra questão é a da rede legada. Uma concessionária que possui várias redes distintas, decorrentes de investimentos realizados em contextos anteriores, não pode se desfazer desses sistemas de imediato, de modo que uma fase de transição será necessária, mas terá como consequência o impacto na qualidade do serviço, na sua gerência e estrutura de custos.

Algumas aplicações são mais exigentes que outras quanto aos requisitos de QoS e são dois os tipos básicos:

- Reserva de recurso – os recursos são alocados de acordo com a QoS da aplicação e submetido à política de gerenciamento.

- Priorização – o tráfego da rede é classificado e os recursos são alocados de acordo com os critérios da política de gerenciamento. Os elementos da rede dão tratamento preferencial para as classificações que tenham mais requisito de banda.

Os tipos de QoS podem ser aplicados a fluxos individuais da aplicação ou para fluxos agregados e existem duas formas de caracterizar: por fluxo, definido um *stream* de dado individual e por agregado, que é caracterizado por um ou mais fluxos, e podem ser associados conforme a natureza de cada serviço, conforme Figura 29.

TIPOS DE TRÁFEGOS				
	Voz	Vídeo	Dados <i>Best Effort</i>	Dados Críticos
Banda (<i>Throughput</i>)	Baixa a Moderada	Moderada a Alta	Moderada Alta	Baixa a Moderada
Sensibilidade à Perda de Pacotes	Baixa	Baixa	Alta	Moderada a Alta
Sensibilidade ao Atraso	Alta	Alta	Baixa	Moderada a Alta
Sensibilidade ao Jitter	Alta	Alta	Baixa	Baixa a Moderada

Fonte: www.cpqd.com.br

Figura 29. Tipos de tráfegos

Definindo SLA (*Service Level Agreement*) como um acordo em nível de serviço é um componente importante do contrato de serviço e especifica o acordo de conectividade e desempenho. A solução de gerência necessita prover meio de gerenciar diversos acordos e permitir ao usuário monitorar individualmente.

Independente da tecnologia ou do protocolo, a rede impõe certas restrições sobre a passagem do tráfego. Uma notável restrição é o impacto no desempenho nas aplicações de rede como o *throughput*, *delay*, *jitter* e *datagram loss*. Diferentes tipos de aplicações podem tolerar cada uma destas restrições de forma diferente [4].

Na implementação da QoS no modelo DiffServ são necessárias as seguintes características (*features*):

- Traffic classification.*
- Queuing e buffer management.*
- Scheduling.*
- Rate-limiting.*

Na classificação do tráfego, ao chegar ao nó *ingress*, os datagramas são classificados em diferentes classes de encaminhamento e da sua prioridade. Os datagramas podem ser marcados de acordo com o campo ToS (*Type of Service*) e os que pertencem a classes diferentes, vão pertencer a filas diferentes, o que irá determinar o *delay* e a latência devido a maior ou menor prioridade do tráfego. A função do escalonador no nó vai decidir a ordem do atendimento nas filas. A limitação da taxa na sua modalidade mais popular é a de policiamento (*policing*) e formatação (*shaping*) do tráfego. Quando o tráfego que chega do cliente excede a taxa negociada, este tráfego será descartado [3].

3.5. ENGENHARIA DE TRÁFEGO

A engenharia de tráfego tem como proposta balancear a carga nos diversos enlaces da rede, evitando que alguns enlaces fiquem congestionados, ou sobrecarregados, enquanto outros estão ociosos, ou pouco utilizados.

Os protocolos de roteamento convencionais, tais como RIP, OSPF e ISIS, baseiam sua decisão de escolha de rota no cálculo de melhor custo (métrica). Porém nem sempre o menor caminho é o que atenderá com eficiência os requisitos de QoS de determinados fluxos de dados. Às vezes, um segundo caminho para determinado destino que, apesar de possuir custo maior, não apresenta congestionamento.

A engenharia de tráfego, além de racionalizar o uso dos circuitos do *backbone*, ainda contém mecanismos de proteção à falha, como a facilidade de estabelecimento de caminhos *backup* caso o caminho principal esteja indisponível. E, os caminhos antes considerados apenas para contingência, agora podem ser utilizados para balanceamento do fluxo de dados entre dois pontos, que contribui para uma melhor ocupação dos circuitos de um *backbone*.

Além disso, a otimização do tráfego em uma rede MPLS-TE, também pode gerar receitas de duas maneiras:

- Oferta de serviços com garantia extra.
- Redução dos custos de investimento em recursos de redes.

Esta redução de custos é atingida pela engenharia de tráfego com o aumento da porcentagem máxima de utilização do *link*. Os operadores constantemente monitoram a utilização do *link* para determinar em que ponto é necessário um agendamento de *upgrade* do *link*. Cada rede tem suas próprias regras de *upgrade*, que podem ser expressas em termos de porcentagem de utilização de *link*, por exemplo, 50% que poderá implicar num *upgrade*. O benefício da engenharia de tráfego neste contexto é que ela permite uma maior utilização dos *links*, porque há um maior controle sobre o caminho que o tráfego segue na rede, ambos sob operação normal e em caso de falha. Melhorando a média de porcentagem de utilização do *link*, seu *upgrade* (expansão) poderá ser postergado [31].

Mas, nem todos os desenvolvimentos em MPLS são usados para engenharia de tráfego e, dependendo do protocolo de distribuição de *label* utilizado, nem todas as redes MPLS podem prover engenharia de tráfego. Portanto, apesar de engenharia de tráfego ser frequentemente confundida com MPLS, uma rede MPLS não é necessariamente uma rede de engenharia de tráfego.

Vale ressaltar que ao discutir os benefícios da engenharia de tráfego, tanto sua implementação quanto sua manutenção devem ser simples para que, em termos financeiros, o custo operacional possa ser justificado pelas receitas trazidas com a engenharia de tráfego. O MPLS provê esta simplicidade operacional requerida, assim como flexibilidade para implementação de políticas complexas de engenharia de tráfego. Ainda hoje em dia não vemos a utilização da engenharia de tráfego de forma automatizada em todas as redes MPLS das operadoras.

3.5.1. FUNCIONAMENTO DA ENGENHARIA DE TRÁFEGO

Prioridade de túneis: alguns túneis são mais importantes que outros, por exemplo, um túnel transportando tráfego VoIP (*Voice over IP*) e outro túnel transportando somente dados, competindo pelos mesmos recursos. Um túnel pode ter seu nível de prioridade definido pelos

valores entre 0 e 7. Quanto maior for o número atribuído ao nível da prioridade, menor a importância do túnel em questão [26].

Cálculo e configuração do caminho: a implementação de MPLS TE pode ser realizada definindo explicitamente o caminho que um LSP deve utilizar. Esta definição pode determinar todos os roteadores intermediários ou apenas parte deles.

Outra forma de implementar MPLS TE é criar LSPs e deixar que a definição dos caminhos seja feita totalmente através de protocolos de roteamento como, por exemplo, o OSPF e o CSPF (*Constraint-Based Shortest Path First* [33]). Estes caminhos podem ser definidos utilizando métricas como custo do enlace, largura de banda, latência, entre outras [1].

Depois que um caminho foi calculado, esse caminho precisa ser sinalizado na rede para estabelecer uma cadeia de salto com os rótulos que representam o caminho e para consumir quaisquer recursos consumíveis (largura de banda) através desse caminho. Essa sinalização é realizada usando RSVP com as extensões para o MPLS TE. O protocolo RSVP, ou o CR-LDP (*Constraint-Based Routing LDP*), permitirão alocar recursos na rede para cada caminho.

Configuração do caminho: após um *headend* completar seu CSPF para determinado túnel, ele terá que sinalizar essa solicitação para a rede, enviando uma mensagem de PATH, este roteador é chamado *upstream* e o que recebe de *downstream*. O roteador *downstream* verifica o formato da mensagem e a quantidade de largura de banda, processo chamado de controle de admissão. Se tiver sucesso, o roteador *downstream* cria uma mensagem PATH e envia ao salto seguinte no ERO e seguem até o último nó ERO, final do túnel MPLS TE.

Quando o final do túnel observa que ele é o destino, responde com mensagem RESV de volta ao *upstream*, que contém uma informação de que a reserva chegou até o final do túnel e contém o rótulo de chegada que o roteador *upstream* deverá usar para enviar pacotes ao longo do LSP TE até o final do período [26].

Métodos para encaminhar tráfego por uma interface de túnel:

- Rotas estáticas.
- Orientado a políticas.
- Rota automática.

Encaminhamento por túnel usando rota automática é o método mais fácil, o administrador configura a rota apontando uma interface do túnel. Encaminhamento por túneis com roteamento orientado a políticas (PBR) é um recurso de encaminhar o tráfego por uma interface específica baseada na política configurada e não na tabela de roteamento.

Encaminhamento por túneis com rota automática é o modo mais detalhado, além do switch é necessário sub-interface ou túnel (GRE - *Generic Routing Encapsulation*) e ativar o IGP nessa interface para formar a adjacência do protocolo de roteamento, descobrir rotas e construir uma tabela de roteamento envolvendo essa interface, sem isso, não poderá descobrir dinamicamente que o tráfego deve ser enviado por essa interface.

Uma das propriedades do MPLS TE é que não é preciso rodar um IGP sobre um túnel MPLS TE, porque os túneis TE são unidirecionais, é difícil rodar IGP em cima dele e como já se tem a topologia do estado de enlace, não precisa inundar outra cópia da informação do roteador da área por um túnel.

3.6. PROTEÇÃO E RESTAURAÇÃO

Um alto grau de disponibilidade de serviço está se tornando imperativo, tanto para o usuário como para o provedor do serviço. Autoregeneração implica em redundância de facilidades e necessidade de inteligência na rede para interpretar rapidamente os resultados dos diagnósticos da rede. Técnicas de autoregeneração que fornecem um alto grau de proteção e uma rápida resposta de restabelecimento, usando o mínimo de facilidades disponíveis.

Segundo Furtado [11] a redução dos custos de proteção de rede mantendo um nível aceitável de sobrevivência se tornou um desafio para os engenheiros e planejadores das redes de telecomunicações. O objetivo é permitir a autoregeneração em todas as partes das futuras redes com uma boa relação custo x benefício [61]. Nos últimos anos um grande número de técnicas de proteção e restabelecimento foi desenvolvido para suportar os atuais e futuros requisitos de sobrevivência das redes de faixa larga. Redundância de componentes, diversidade de rotas, anéis autoregenerativos e restabelecimento dinâmico estão entre as soluções para auxiliar partes específicas de redes a se recuperar das falhas de um ou mais elementos da rede. Algumas dessas técnicas ainda estão em estudo, enquanto outras já foram amplamente discutidas e inclusive já estão disponíveis a planejadores e operadores de rede,

segundo a RFC3386, aplicáveis as redes ópticas SDH / ASON / MPLS. Seguem alguns dos principais tipos de proteção de caminho:

Proteção SNC-P (*Subnetwork Connection Protection* - [61]) - é um mecanismo de proteção dedicada que pode ser usada em qualquer estrutura física de rede, malha, anel, estrela ou uma topologia mista. Este tipo de proteção é muito utilizado em topologias de redes com anéis unidirecionais. Este mecanismo de proteção pode ser usado para proteger um trecho do caminho (por exemplo, a um trecho onde dois segmentos de caminho separados estão disponíveis). Este mecanismo permite a comutação de falhas da camada servidora (usando monitoração inerente) ou usa a informação da camada cliente (usando monitoração não intrusiva) [11].

Uma proteção SNC (*SubNetwork Connection* - [61]) de serviço, em caso de falha ou degradação de desempenho, é substituída por uma SNC de proteção de via. Uma SNC corresponde a um segmento de via. Esta é uma proteção dedicada 1+1 na qual o tráfego de serviço e o tráfego de proteção no ponto de transmissão da conexão de sub-rede são transmitidos em dois caminhos separados. No caso da proteção dedicada 1+1, o tráfego é transmitido tanto nas conexões de sub-rede (SNC) de proteção quanto de serviço. No lado de recepção da SNC, a comutação de proteção é efetuada através da seleção de um dos sinais baseados apenas na informação local. A proteção SNC-P genericamente protege contra falhas na camada servidora, e falhas e degradações na camada cliente. Pode operar no modo reversível ou não reversível. O tempo de comutação do algoritmo para SNC-P deve operar de forma tão rápida quanto possível. Um valor de 3 ms foi proposto como tempo objetivo, onde também ocorre a duplicação de banda. Como um princípio geral, para cada direção de transmissão, o canal de proteção deve seguir um roteamento separado do canal de serviço [11].

A proteção MS-Spring (*Multiplex Section Shared Protection Ring* - [61]) - os anéis que operam com proteção compartilhada de seção de multiplexação, são denominados MS-Spring para anéis com proteção de seção compartilhada, os canais principais transportam os serviços a serem protegidos enquanto os canais de proteção são reservados para a proteção deste serviço. O tráfego de serviço é transportado bidirecionalmente sobre os arcos do anel, um sinal de tributário entrando em um nó trafega em uma direção enquanto o seu tributário de saída associado trafega no sentido oposto, mas sobre o mesmo arco. O par de tributários usa apenas a capacidade ao longo dos arcos entre os nós onde o par é derivado e inserido. Então, o

padrão segundo o qual esses pares de tributários são colocados no anel tem impacto sobre a máxima carga que pode ser transportada nos anéis MS-Spring. A soma dos tributários que atravessa determinado arco, não pode exceder a máxima capacidade deste arco [11].

Uma vantagem dos anéis MS-Spring é que o serviço pode ser roteado no anel tanto numa direção quanto na outra, através o caminho longo (*long path*) ou o caminho curto (*short path*). Embora o caminho curto seja usualmente o preferido, ocasionalmente o roteamento dos serviços através do caminho longo permite certo balanceamento do tráfego no anel, otimizando assim sua capacidade.

Quando os canais de proteção não estão sendo usados para restaurar os canais de serviço, estes poderão ser usados para transportar tráfego extra. No caso de um evento de proteção por comutação de anel, o tráfego nos canais principais será transferido para os canais de proteção causando a remoção do tráfego extra dos canais de proteção. Ou, seja não se pode assegurar a proteção ou continuidade do tráfego extra, que não tem nenhuma prioridade. O tráfego extra poderá ser entendido como dados não prioritários, pois em caso de qualquer falha na rede este tráfego pode ser interrompido.

Durante a comutação do anel, os canais transmitidos através do arco em falha são transferidos em um nó de comutação para os canais de proteção e transmitidos na direção oposta à ocorrência da falha. Este tráfego viaja através do caminho longo do anel nos canais de proteção para o outro nó de comutação onde os canais de proteção são transferidos de volta para os canais principais. Na outra direção, os canais de serviço são comutados de maneira similar ao descrito acima. Todos os anéis com proteção de seção de multiplexação compartilhada suportam comutação de anel.

Cada fibra transporta tanto os canais de serviço quanto os canais de proteção. A metade dos canais é definida como prioritários e metade são definidos como canais de proteção. Os canais principais são protegidos pelos canais de proteção que trafegam na direção oposta ao longo do anel. Isto permite o transporte bidirecional do tráfego de serviço. Neste tipo de proteção temos o efeito “trombone”, que é o efeito de percorrer todo caminho até aonde ocorreu a ruptura e depois voltar e percorrer todo caminho de volta até alcançar o nó de destino. O tempo de comutação, a proteção MS-Spring deve operar de forma tão rápida quanto possível. Um valor de 50 ms foi proposto como tempo objetivo. Não há duplicação da banda para proteção do tráfego de serviço, como ocorre na proteção SNC-P. No SDH o

protocolo K1/K2 é o protocolo que comanda a comutação do tráfego de serviço na proteção MS-Spring.

Proteção MS-Spring Transoceânico [61] - A proteção MS-Spring Transoceânico atua em sistemas de enlaces de longas distâncias em cabos ópticos submarinos intercontinentais, onde pode ocorrer uma falha entre os nós do anel. Quando ocorre a falha, a proteção executa a função de comutação entre os nós afetados pela falha, não havendo a necessidade do sinal percorrer todo o caminho até as estações adjacentes a falha para o sistema ser restaurado. Este tipo de proteção é utilizada em enlaces intercontinentais, com isso estamos evitando o efeito “trombone” que seria prejudicial para a restauração, pois o tempo de convergência seria maior. O tempo de comutação para a proteção MS-Spring Transoceânico deve operar de forma tão rápida quanto possível. Um valor de 10 ms foi proposto como tempo objetivo. Não há duplicação da banda para proteção do tráfego de serviço. O protocolo *Extended K1/K2* é o protocolo que comanda a comutação do tráfego de serviço na proteção MS-Spring Transoceânico [11].

-Restauração:

A restauração tem um papel importante nas telecomunicações, pois nos dias atuais as redes são conhecidas como redes multiserviços. De uma forma genérica, esta é uma funcionalidade que pode ser agregada a rede de forma a manter a conexão dos usuários, para que as informações não sofram perdas. Nos casos de perda da conexão as mesmas tendem a ser restauradas, conforme RFC3469. Seguem a seguir alguns tipos de restauração:

-Unprotected:

O circuito não será protegido contra falhas (neste caso não existe proteção em caso de perda de conexão). Este tipo de configuração é feita em circuitos de baixíssima prioridade.

Seguem alguns dos principais tipos de proteção de circuitos:

-PRC (*Protection Restoration Combined*).

O circuito de proteção é previamente alocado (tempo de restauração em até 50 ms). É configurado para circuitos com altíssima prioridade (recuperação ultrarrápida). É criado um circuito com a Proteção SNC-P. Após a falha do circuito principal, o tráfego é comutado para o circuito de proteção e o mecanismo de restauração cria um circuito SNC-P para proteção do tráfego, garantindo assim que o circuito continue totalmente protegido.

-Vantagem:

- Garante a recuperação em um tempo menor que 50 ms.
- Uso flexível e eficiente de recursos de proteção da rede.
- Recuperação ultrarrápida das falhas sem intervenção manual.

-GR (*Guaranteed Restoration*)

O circuito de proteção é previamente calculado, mas não está alocado (restauração em até 100 ms). Para cada circuito principal é calculado um circuito de proteção que pode ser compartilhado com a proteção de outros circuitos, o mecanismo da restauração ativa o circuito de proteção após a falha, e um circuito de proteção será calculado após a falha, garantindo assim que o circuito continue totalmente protegido.

-Vantagem:

- Recuperação do circuito principal para toda a falha na rede.
- Uso flexível e eficiente de recursos de proteção da rede.
- A recuperação do circuito de proteção é feita usando a proteção SBR (não reserva banda).
- Atua na recuperação de falhas múltiplas sem intervenção manual.

Dois circuitos podem compartilhar a mesma via de proteção. Caso haja um rompimento simultâneo nos circuitos que estão compartilhando a via de proteção, o circuito que será protegido será aquele que for mais prioritário.

-SBR (*Source Based Reroute*):

O circuito de proteção é calculado no momento em que ocorre a falha no circuito principal (tempo de recuperação maior que o GR). A comutação ocorrerá através de mecanismos de prioridade, o circuito de proteção usará caminhos disponíveis do circuito principal quando for possível.

-Vantagem:

- Garante a recuperação do circuito para toda a falha da rede.
- Uso flexível e extremamente eficiente de circuitos de proteção da rede.
- Recuperação de falhas múltiplas sem intervenção manual.

-No caso de o GR e o PRC não encontrem o circuito de proteção, o SBR é usado para salvar o tráfego.

SNC-P (*SubNetwork Connection Protection*)

Proteção para uma única falha (restauração em até 50 ms), ela utiliza banda alocada para proteção tendo um alto custo para a Operadora e para os usuários. Ela ainda é usada para o processo de reversão aplicado na recuperação das falhas.

-Reversão:

A reversão é feita quando o circuito principal está disponível outra vez, para a reversão é criado um circuito SNC-P provisório. Quando o circuito principal estiver livre de falha ou erro, o circuito SNC-P retorna para o circuito principal. A verificação do circuito principal dura de 2 a 6 minutos. O SNC-P é removido depois do circuito proteção para o circuito principal.

-Vantagem:

- Procedimento automatizado para reverter o tráfego ao seu leito original.

- O uso do SNC-P impede sucessivas comutações durante o restabelecimento do circuito principal ao seu leito original.

- Evita a fragmentação da rede.

Quando falamos em proteção, é na configuração dada a rede para quando ocorrer falhas na rede, a rede deve estar capacitada para tomar uma decisão, ou seja, é anterior a falha. A restauração é a reação da rede a um evento de falha, que vai seguir o esquema pré-definido na proteção estabelecida [32].

3.7. EVOLUÇÃO DO PERFIL DE TRÁFEGO E REDE

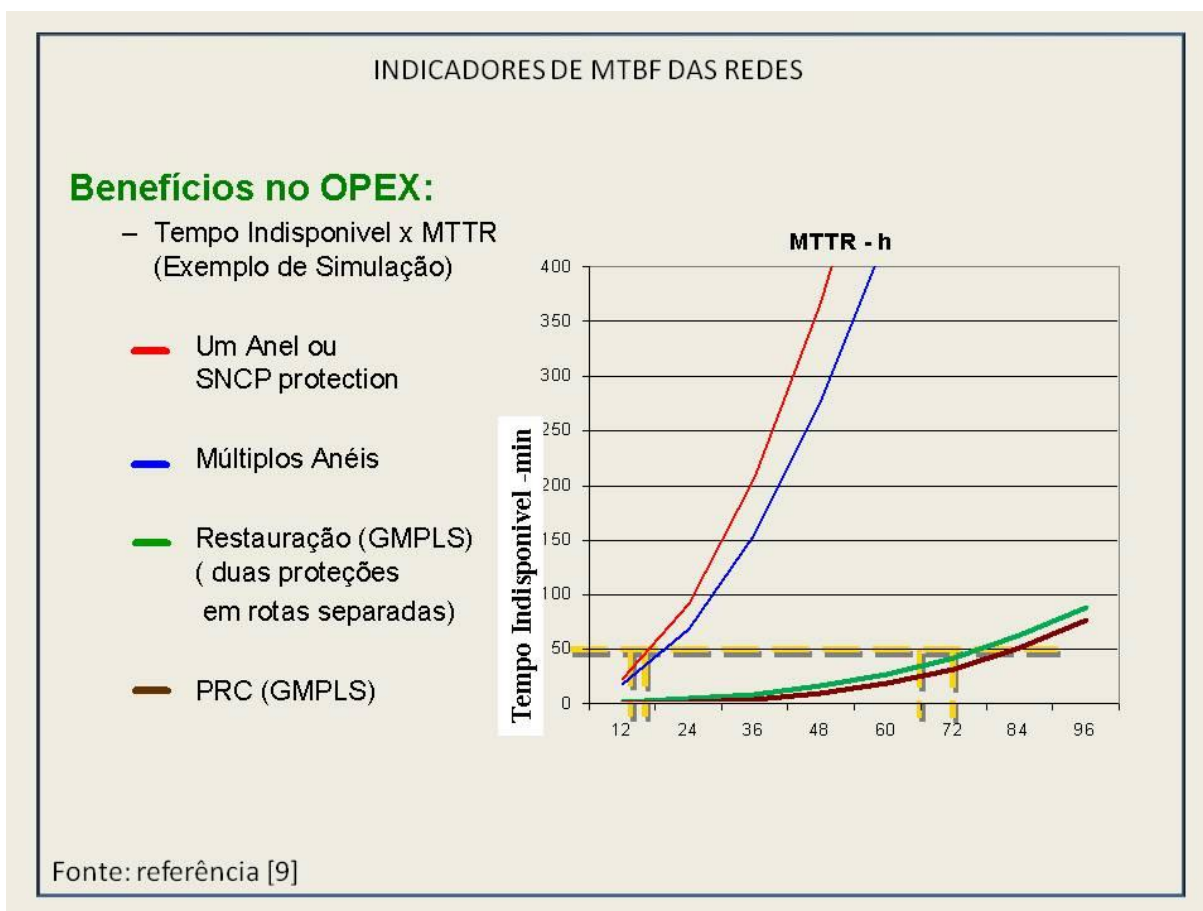


Figura 30. MTBF das Redes

Outro parâmetro importante a ser analisado, através da Figura 30, é a análise do gráfico *down-time* (tempo em que a rede fica fora, um pedaço da rede fica fora, ou que um *link* da rede fica fora do ar), comparado com o tempo necessário para recuperar a rede, dependendo do *down-time* da rede, é possível prolongar tempos de reparo, consequentemente, diminuindo o OPEX da rede.

Os benefícios na restauração utilizando GMPLS podem ser entendidos mais claramente, onde é verificado que quanto maior a disponibilidade da rede, maiores alternativas de caminho para restaurar o tráfego em caso de perda. Isto significa redução do OPEX, através da menor necessidade de pessoal para realizar a manutenção da rede e menos unidades sobressalentes. Consequentemente, o dimensionamento de recursos, principalmente de reparo, é menor, permitindo mais tempo para recuperar uma possível falha e possibilitando uma redução da despesa operacional da rede.

Resumindo as proteções podem ser de rede ou de equipamento, as proteções de equipamento podem ser aplicadas com placas *backup*, como matriz, fonte, *fan*, etc..

As proteções de rede podem ser:

- DWDM
 - Proteção de linha.
 - Proteção de canal.
 - Proteção SNCP.
 - Proteção ODU.
 - Proteção OWSP.
- SDH
 - Proteção MSP (1+1 ou 1:N).
 - Proteção MSP transoceânico.
 - Proteção SNCP.
 - Proteção SNCTP (túnel).
- ETHERNET
 - Proteção LCAS.
 - Proteção LAG.
 - Proteção STP / RSTP.
- ASON
 - Proteção combinada.
 - Proteção garantida.
 - *Reroute*.
 - Proteção SNCP.

3.8. RESUMO DO CAPÍTULO

Diversas redes acadêmicas realizaram estudos para viabilizar uma interoperabilidade global, com sistemas que realizassem reserva e provisionamento em redes híbridas. O estudo colaborativo entre as NRENs desenvolveu as ferramentas DRAGON e OSCARS, além de novos protocolos interredes, para a oferta de serviços dinâmicos (*e-science*, por exemplo) em redes DCNs.

Para resolver os problemas de interoperabilidade das redes de pacotes, o ITU-T e o IETF se uniram para formar o *Joint Working Team* (JWT) e desenvolver a tecnologia MPLS (híbrida) e os padrões relacionados. Esse grupo criou o *MPLS Transport Profile* (MPLS-TP), que absorve tecnologias de camada de transporte para OAM, proteção e gerência. Os trabalhos deste grupo incluem arquitetura, características funcionais, aspectos de estrutura e mapeamento, aspectos de gerenciamento e camada física.

Para oferecer serviços *e-science*, temos que pensar nos requisitos de crescimento das redes com garantias, como: QoS nas camadas de transporte e aplicação, redução dos custos de rede, padronização das funcionalidades de controle da camada física, criação de circuitos fim a fim e de forma dinâmica, etc.

Com base nos estudos das tecnologias atuais (Capítulo 2), das redes híbridas, QoS e de técnicas de restauração e proteção deste Capítulo 3, apresentaremos como oferecer o serviço Profile Service com a especificação da interface, primitivas e modelos.

4.PROPOSTA DE NOVO SERVIÇO

Em 2002 nasceu o conceito da rede híbrida, onde uma mesma infraestrutura de comunicação é usada para prestar serviço de pacotes IPs e para circuitos fim a fim para aplicações selecionadas. Desde então, a opção pela rede híbrida tem sido adotada por todas as redes de pesquisa, que oferecem serviços de grande capacidade (*Gigabits*). A rede do futuro deve ser então, uma rede híbrida, através da qual a rede manterá os atuais serviços de conectividade e adicionará o serviço de comutação dinâmica de circuitos de camada 2. O desempenho do serviço na rede reflete a composição dos desempenhos da conexão de acesso, do *backbone* de transporte e do próprio servidor que executa o *software* da aplicação.

Foi realizado estudo de capacidade da rede, bem como suas vulnerabilidades de segurança. Outros estudos estão em desenvolvimento, para obter escalabilidade, confiabilidade e disponibilidade da plataforma e do serviço. As redes DCN abrem grandes e inovadoras perspectivas para o mercado de telecomunicações permitindo atender a demandas reprimidas e desenhar novos serviços. No escopo de uma NREN, o serviço DCN permite suportar e mesmo fomentar novas atividades de *e-science*, computação em *grid* e trabalhos colaborativos.

A contribuição desta dissertação, é a de especificar as funcionalidades e as necessidades de adaptação de sistemas e processos para entrega do serviço Profile Service nas redes dos clientes, implementando de forma transparente, um serviço de comunicação de dados que utiliza a arquitetura das redes híbridas e comutação dinâmica para prover circuitos ponto a ponto de altas capacidades, com a solicitação a partir do cliente, com proposta de forma organizada com um plano de planejamento e de implementação, descrição das primitivas do serviço e políticas, que organizam as lógicas do serviço, topologia e QoS.

Este serviço foi planejado com 3 modalidades, que possuem um conjunto de atributos que definem as suas características (Gold, Silver L2 e Best Effort) e viabilizado para 3 cenários, o atendimento ao cliente que demande um dos serviços *e-science*, considerando o cliente como um domínio particular interagindo com o domínio da rede do provedor, em um segundo cenário o cliente possuindo seu próprio domínio DCN e no terceiro cenário o cliente participando do domínio DCN do provedor.

Para os três cenários descritos podem ser reavaliados testes dos serviços, quando as novas placas dos equipamentos, possuírem a interface UNI 2.0 implementada, então poderemos oferecer os serviços dinâmicos de forma automática. Os serviços são mapeados com definição de QoS em uma fila específica nos roteadores do provedor, de forma a garantir que o tráfego da rede DCN, fique separado do tráfego da rede que está em operação.

Este capítulo apresenta as especificações deste serviço para as aplicações *e-science* e sua implementação em redes híbridas.

4.1. APLICAÇÕES

Segundo Spadaro[55], os acessos de banda larga estão se tornando cada vez mais uma exigência no mundo, tendo em vista a crescente demanda por serviços de voz, dados, vídeo, *Internet* e multimídia, atividades *e-science*, computação em *grid* (combinação de recursos de computação com múltiplos domínios administrativos) e trabalhos acadêmicos colaborativos.

O termo *e-science* é usado para aplicações do mundo acadêmico que transferem grande volume de dados e no Brasil é chamado de “e-infraestrutura” (wikipédia), geralmente usado na computação científica, com recursos computacionais como *cluster* (central e secundário) e ainda com servidores de armazenamento. Um exemplo dessa infraestrutura é a Redeh da RNP, citado no item 3.2, ou seja, uma infraestrutura que permite a distribuição dinâmica de conteúdos como se fossem sítios na *Internet*.

O grande desafio é a criação de ferramentas que alterem a configuração da rede dinamicamente de forma programada para atender aos requisitos de transferência dos volumes dos dados com banda ajustada (a estas necessidades) e com confiabilidade. Como propostas surgem: 1) o incentivo a produção científica, levando a geração de novas interfaces *web*; 2) gerenciamento de projetos de pesquisa e 3) serviços que permitam ao usuário escolher a melhor aplicação.

Em setembro de 2012, ocorreu o *workshop* EUBrazil / OpenBio em Recife, sobre infraestrutura para pesquisa científica, projeto colaborativo com o Ministério da Ciência e TI (CNPq) e a Comissão Europeia, com o objetivo de implementar acesso aberto para a comunidade científica, com a participação da RNP [68].

A contribuição da presente dissertação é a proposta de criação de um serviço que disponibilize no *website* do cliente, alternativas de acessos, onde ele vai realizar o seu *login* e mapear cada aplicação através da URL (*Uniform Resource Locator*) do provedor. O projeto DRAGON (item 3.2.2) estuda as questões de estabelecimento das conexões intradomínio. Com base nos estudos acadêmicos, foi desenvolvido o protocolo IDC, para dinamicamente aprovisionar recursos no domínio da rede (banda, VLAN, etc.) [68].

Para esta aplicação, ao implementar um protocolo bem definido, com característica cliente - servidor, poderá interagir, usando um número de porta associado, conforme Figura 31.

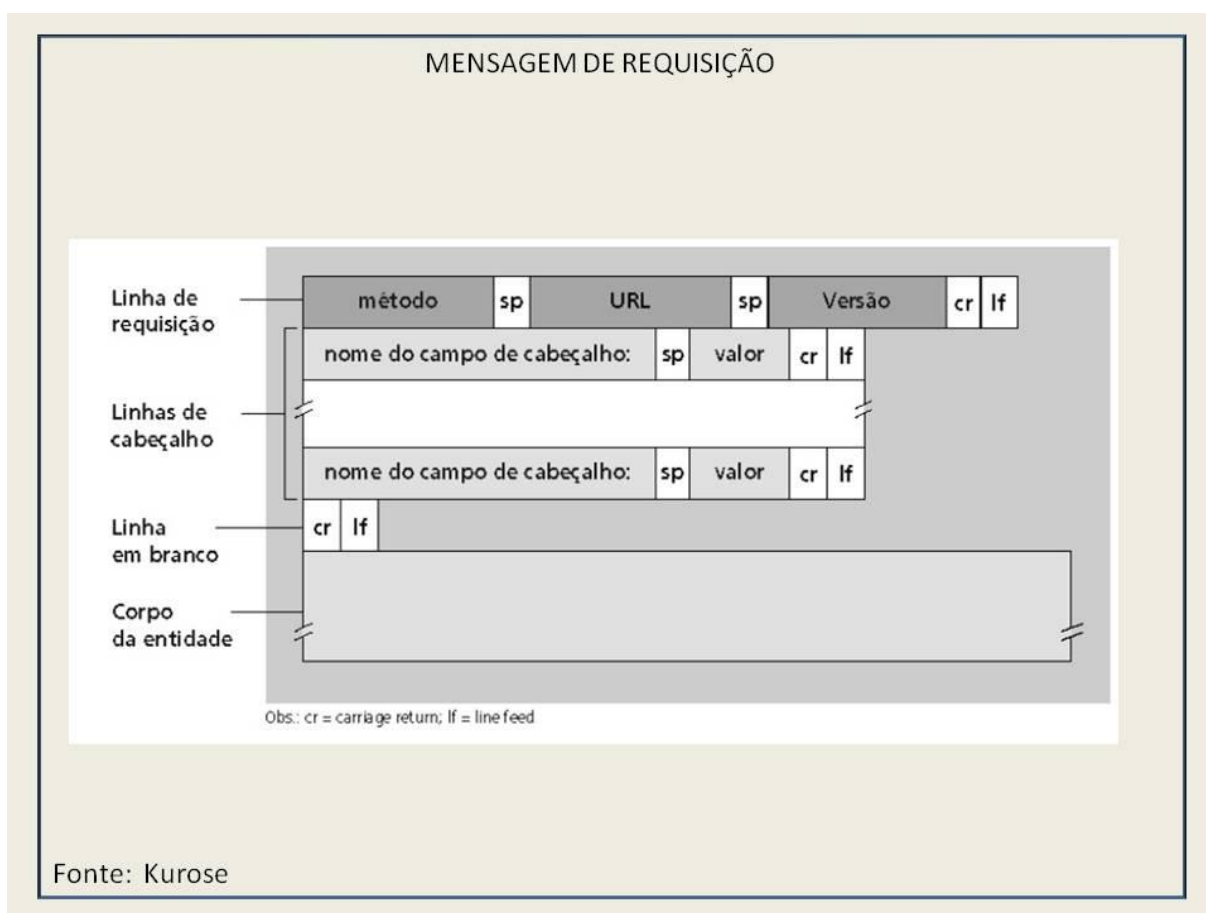


Figura 31. Formato Geral de Mensagem

4.2. DETALHAMENTO DO SERVIÇO

Podemos definir um serviço pela forma como ele é visualizado pelo cliente (CE), que inclui as interações entre as interfaces UNI do cliente e da rede, a esta interação chamamos de

conectividade ou EVC (Conexão Virtual *Ethernet*), cujos atributos são UNI ID, MAC e CE-VLAN ID, estes serviços podem ser: ponto a ponto, ponto-multi-ponto ou uma combinação destes.

Como foi visto no item 2.9, o serviço E-LINE é o mais simples fornecido pela rede. Ele se constitui de uma ligação ponto a ponto entre 2 nós da rede. O serviço E-LAN é um serviço com ligações multiponto entre os nós da rede e o serviço E-TRE são ligações ponto - multiponto, não deixando de ser um caso particular do serviço E-LAN. Os serviços são prestados nas camadas L2 e com QoS garantido pelo campo do protocolo IEEE 802.1p [19].

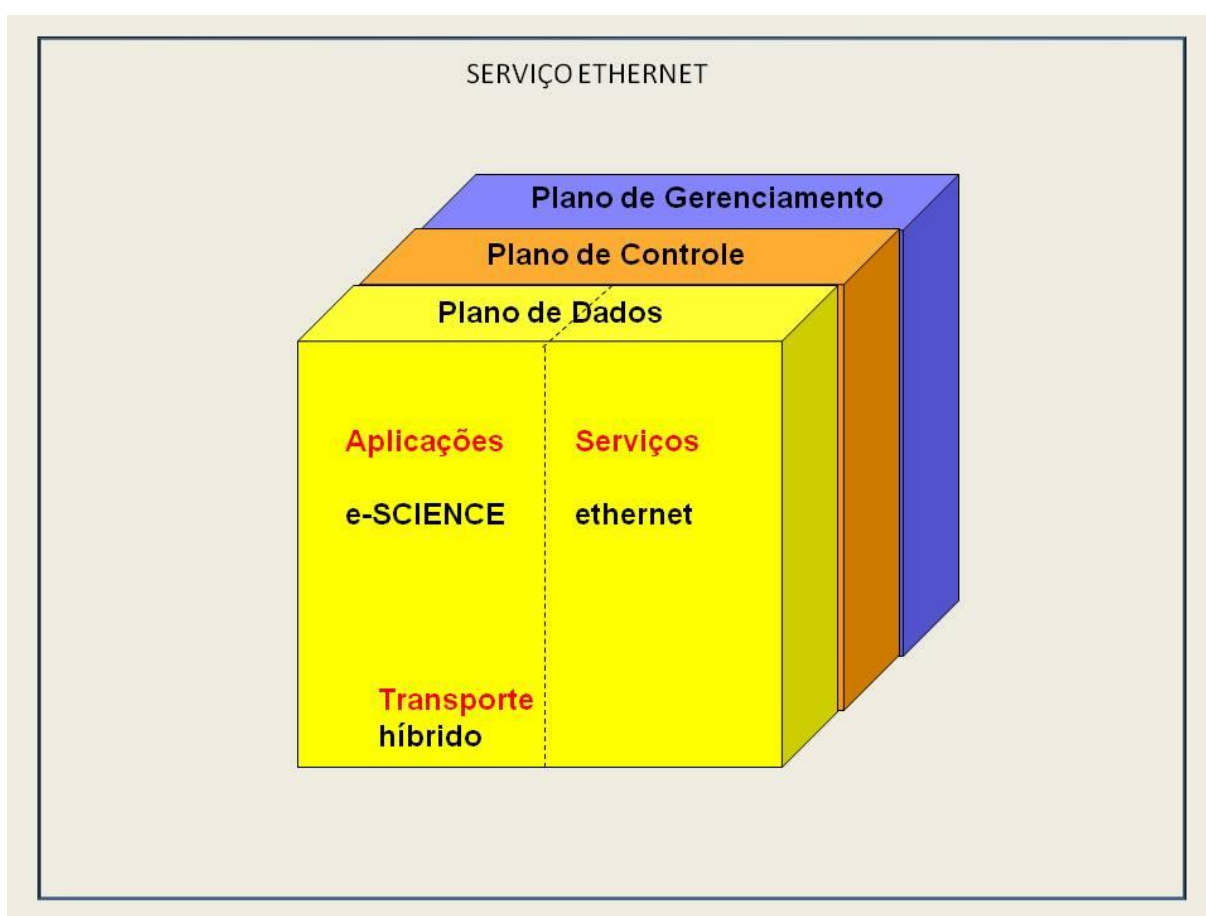


Figura 32. Modelo da Rede

Os serviços vão demandar os seguintes requisitos para aplicação:

- Identificação do usuário - nome e senha.
- Análise e negociação - prioridade.
- Especificação - atende ao solicitado.

- Validação - consistência dos dados.

Há ainda a necessidade de novos procedimentos de pedido e de estabelecimento de conectividade dinâmica entre os diversos dispositivos do cliente e a rede de transporte. O desenvolvimento de tais procedimentos requererem a definição das interfaces físicas do cliente e da rede, os serviços, protocolos de sinalização, mecanismos de mensagens e procedimentos de autodescoberta, conforme apresentado na Figura 32. Estas definições constituem a interface de rede do usuário (UNI), descrita conforme o padrão de interface definido pelo *Forum* OIF-UNI-01.0, que permitiu implementar uma camada da rede óptica de transporte aberta e com conectividade dinâmica para as camadas do cliente (IP, ATM, SDH, etc.), o que possibilita a implementação da nossa proposta [52].

O padrão utiliza as normas e especificações disponíveis das demais organizações como ITU-T [59], IETF [32] e ANSI (TIA/EIA/568). Para o serviço Profile Service, usa as especificações baseadas no ITU-T (SDH), para os protocolos de sinalização e autodescoberta utiliza os trabalhos do IETF GMPLS, o protocolo LMP e ainda os protocolos de sinalização LDP e RSVP-TE. As atualizações destes protocolos são incorporadas a UNI pela OIF.

O principal serviço oferecido sobre esta interface é a capacidade de criar e excluir uma conexão, sendo a conexão um circuito de banda fixa entre a entrada e a saída do ponto de acesso. Basicamente, existem duas maneiras de iniciar uma chamada de serviço: uma chamada direta e uma indireta. Em ambos os modelos o lado do cliente e o lado da rede são referidos como UNI-C e UNI-N. O cliente solicita o serviço diretamente sobre a interface e, no modelo indireto, uma entidade *proxy* UNI-C desempenha as funções em nome do cliente, ou seja, o *proxy* realiza o serviço de descoberta, resolução de endereço e sinaliza para a UNI-C. Entre a UNI-C e a UNI-N a configuração será, pré-determinada inicialmente, até que a nova versão da interface UNI 2.0 esteja comercial [22].

Atualmente, algumas considerações empresariais e operacionais (como políticas de segurança) levam a uma arquitetura de plano óptico de controle e protocolos que habilitam práticas de negócios, como gestão de acordos por domínios ou políticas. A recomendação ITU-T G.8080 (ASON) pode ser utilizada para expressar essas diferentes responsabilidades administrativas, como endereçamento, restauração, plano de controle, etc. (Figura 32). Os domínios são estabelecidos pelas políticas do operador da rede, que representa um conjunto de entidades agrupadas para uma finalidade específica que prevê o encapsulamento das informações e a caracterização do domínio.

Um domínio, portanto, representa um conjunto de entidades que são agrupadas para um propósito particular, sendo o controle de um domínio composto pela coleção dos componentes do plano de controle.

4.3. SERVIÇO UFF

Como proposta de novo serviço, temos a prestação de serviços estabelecidos pelo cliente em três modalidades e com agendamento automático através das redes híbridas, que trazem a grande vantagem de permitir o estabelecimento e a remoção de circuitos de forma automática. É possível também agendar a utilização de circuitos em períodos pré-determinados, ganhando em flexibilidade e maximizando a utilização dos recursos da rede. Outro ganho importante é a simplificação da operação da rede, retirando a tarefa de estabelecer e remover os circuitos de mãos humanas. Assim, com estas características, várias limitações tecnológicas são derrubadas e um novo horizonte de possíveis serviços se abre. A demanda cada vez maior do mercado por convergência, qualidade, alta capacidade e agilidade no provimento de meios de comunicação, pode ser atendida com o desenho de serviços, usando uma infraestrutura de rede híbrida (item 3.1.1) com comutação dinâmica.

4.3.1. REQUISITOS DE TOPOLOGIA

Nossa aplicação pode ser categorizada dentro dos serviços *e-science*. Esta nossa aplicação batizamos de Profile Service. Trata-se de um serviço de telecomunicações para transmissão de dados baseados no protocolo de enlace *ethernet*. Este serviço se aplica fundamentalmente a interconexão de redes locais de clientes, de forma transparente a protocolos de camadas superiores.

Este serviço tem três modalidades:

- Ponto a ponto: Conectividade entre dois *sites* situados em endereços diferentes seja na mesma localidade ou em localidades diferentes, através de uma conexão ponto a ponto.

- Ponto-Multiponto (*Hub and Spoke*): Conectividade entre um site central e dois ou mais sites remotos, sejam eles na mesma localidade ou em localidades diferentes. Com

este serviço, um site remoto pode se comunicar com outro site remoto por meio do site central.

- Multiponto-Multiponto (*Full Mesh*): Conectividade entre todos os demais *sites* remotos sejam eles na mesma localidade ou em localidades diferentes. Este serviço oferece melhor disponibilidade, mas pode causar atrasos e custos elevados.

4.3.2. CLASSES DE SERVIÇO

Com relação ao QoS, as modalidades acima devem ser divididas em duas categorias e que são detalhados neste item:

- Sem QoS – Não haverá tratamento específico para os pacotes trafegados
- Com QoS – Haverá a possibilidade de criação de 3 classes de serviço para priorização de tráfego na rede do cliente em fila exclusiva

Conforme detalhado na Figura 33, o padrão 802.1p permite que as UNIs da rede de acesso sejam mantidas através do *backbone* (tunelamento), este tunelamento é feito através de duplo *tag* de 4 *bytes*. O mapeamento do CoS ID será de acordo com a prioridade do tráfego, variando de 0 a 7. A Figura 32 também correlaciona as classes de serviço definidas no nosso modelo.

CLASSES DE SERVIÇOS – 802.1p Profile Service				
Classe de Serviço	Aplicações	CoS ID	EVC por CoS ID	Performance serviço
Gold	Videoconferencia	6 , 7	CIR > 0 EIR = 0	Delay < 5 ms Jitter < 1 ms Loss < 0,001%
Silver L2	Transferencia de dados	3 , 4 , 5	CIR > 0 EIR < UNI	Delay < 15 ms Jitter < 1 ms Loss < 0,01%
Best effort	internet	0 , 1 , 2	CIR > 0 EIR = UNI	Delay < 30 ms Jitter < 1 ms Loss < 0,5%

Figura 33. Classes de Serviços 802.1p

4.3.3. CONECTIVIDADE DE REDE

Deve haver conectividade entre todos os *sites* do cliente, na mesma localidade ou em localidades diferentes, permitindo que todos os *sites* tenham capacidade de se comunicar entre si.

Conforme resumido na Figura 34 o serviço está disponível nas velocidades de 100Mbps e 1Gbps. Exceto na modalidade Ponto a ponto na qual a velocidade do acesso é igual nas duas pontas, às velocidades dos acessos nas outras modalidades poderão variar, formando combinações a partir das velocidades disponíveis. O provisionamento da rede para formar circuitos com a rede *ethernet* é de forma transparente.

Apesar de ser transparente aos protocolos do cliente, é possível a entrega de serviços distintos em um mesmo acesso, caso sejam definidos novos serviços, segregados por VLANs (redes privadas formadas em nível de enlace *ethernet*, também conhecidos como *pseudowires*), sendo possível aplicar priorizações de tráfego (QoS) distintos por serviço em

toda a rede *ethernet*. Também por meio deste mecanismo o serviço poderá ser segregado por aplicação, em um único acesso (ex.: VLAN1 – Dados, VLAN2 – Voz), com priorizações de tráfego específicas. Considerando que a capacidade de formar VLANs somente pode ser implementada na rede *ethernet*, é importante ressaltar que em circuitos formados em redes distintas (redes híbridas), cada VLAN será mapeada para um EPL distinto, conforme a Figura 34 [51].

SERVIÇO ETHERNET	
Atributos	Parâmetros
identificação da interface UNI	CE VLAN-ID
Velocidades	100Mbps / 1Gbps
Modo	Full duplex
MAC	IEEE 802.3
Numero de ECV	> 1
Agregação	Sim

Figura 34. Parâmetros do Serviço

4.3.4. ATRIBUTO DO SERVIÇO

A rede *ethernet* permite também interagir de forma não transparente (ex: *spanning tree* do cliente com a rede). A interação de protocolos da rede do cliente com a rede do provedor deverá ser um serviço opcional.

São atributos do serviço:

-Interface.

- Perfil de banda.

- Desempenho.

- Qualidade.

As interfaces especificam o modo como será realizada a comunicação entre o cliente e a rede e deverá ser nos padrões 100Base FX ou 1000 Base SX. O perfil da banda será estabelecido de acordo com as 3 modalidades do serviço Profile Service. O desempenho vai depender da prioridade de cada tráfego, a forma como será mapeado CoS ID, o serviço terá determinado *delay*, *jitter* e *loss*.

São atributos do tráfego:

- Identificador.

- Endereço.

- Interface.

O CE-VLAN ID é o identificador do conteúdo de um *frame* de serviço, que permite que este *frame* seja associado a um EVC na UNI. O endereço será seu MAC.

Deverão estar disponíveis para o cliente mecanismos de monitoração de tráfego e desempenho dos circuitos, acessíveis via *web* (HTTP (*Hypertext Transfer Protocol*)) por meio de *login* individualizada por cliente, conforme apresentada na Figura 35.

Administração de Usuários e-science

USUARIO

➔

Nome: (e-mail)
 Senha: (password)

(escolha o serviço desejado)

CÓDIGO	SERVIÇO	MÁQUINA
 1	Graduação	IP 10.210.144.150
 2	Biblioteca	IP 10.210.144.160
 3	e-mail corporativo	IP 10.210.144.170
 4	e-science - projetos	IP 10.210.144.180
 5	e-science - banco dados	IP 10.210.144.190
 6	e-science - aplicativo agendamento	IP 10.210.144.200

Figura 35. Interface de usuário e-science

Para algumas aplicações, como *e-mail*, transferência de arquivos e aplicações financeiras, não pode haver perda de dados, outras aplicações como vídeos armazenados, podem tolerar certa perda. O efeito que estas perdas causam a aplicação depende dos protocolos usados pela rede. Em virtude do grande volume de dados para os serviços *e-science*, é necessária implantação de infraestrutura robusta e adaptável ao modelo, com os recursos alocados de forma dinâmica, para isso a definição dos planos de controle e gerenciamento a solicitação automática sem intervenção do administrador e ainda das políticas de autenticação.

4.3.5. SINTAXE DA APLICAÇÃO PROFILE SERVICE

As aplicações trocam mensagens com seus pares em outras máquinas.

O serviço Profile Service usa o protocolo HTTP, através de um programa para o cliente com interface *web*. Quando o usuário abre a URL do serviço, seu navegador abre uma conexão TCP com o servidor na máquina principal do serviço *e-science*, que exibe os arquivos em HTML (*HyperText Markup Language*). Eles incluem URLs que apontam para outros arquivos e seu navegador *web* reconhece e solicita abertura de URLs, que são os *links* hipertexto.

Quando o usuário acessa a página do servidor usando HTTP, este usa o protocolo TCP. Nesta página ele solicita que seja criado um perfil dentro do portal da Universidade que ele estuda ou trabalha.

Basicamente cada mensagem HTTP possui o formato geral:

- Linha_inicial.
- Cabeçalho_mensagem.
- Corpo_mensagem.

A linha inicial indica se é uma mensagem de solicitação ou de resposta. O cabeçalho especifica opções e parâmetros que qualificam a solicitação.

Mas para que o usuário estabeleça uma conexão, ela utiliza de primitivas de serviço é que a camada superior inicia uma transferência de dados com a camada correspondente do usuário remoto e estabelece a chamada para o serviço, diretamente na sua interface de cliente.

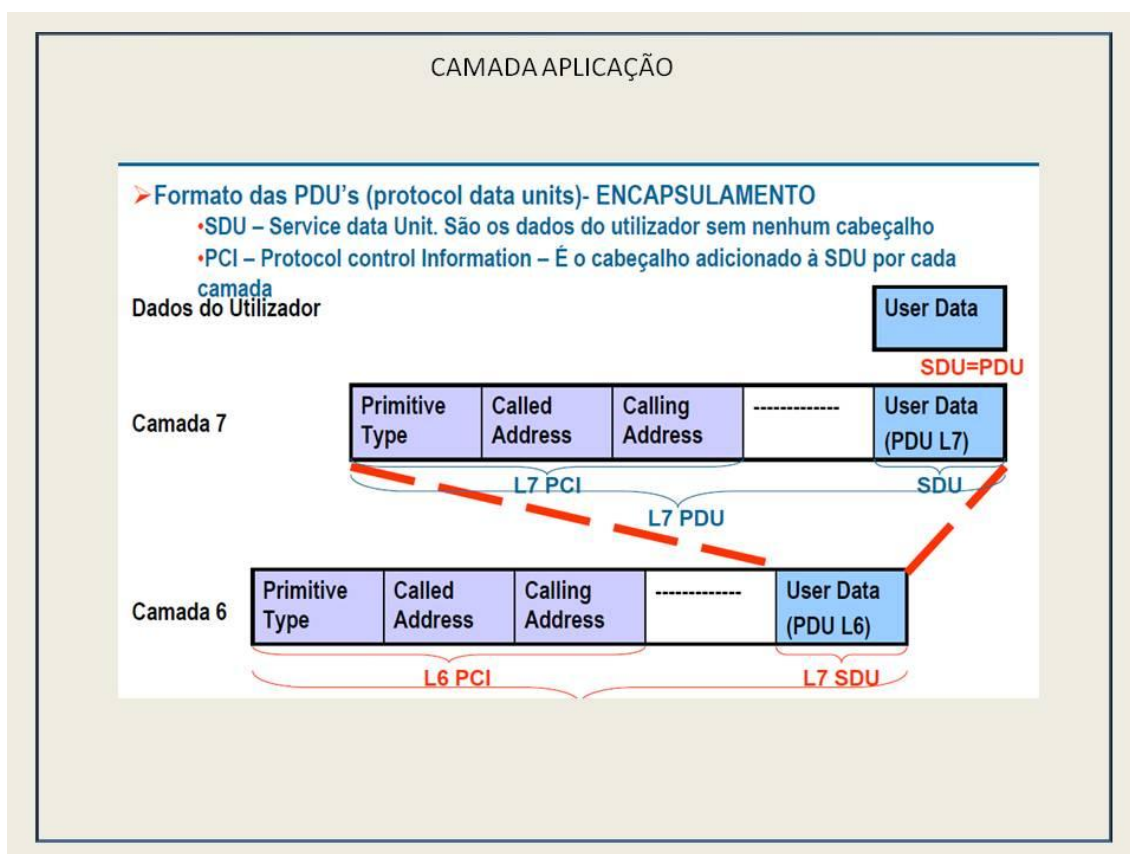


Figura 36. Formato da PDU

Conforme verifica-se na **Figura 36**, o encaminhamento da camada de aplicação é através das PDUs L7 (*Protocol Data Units*) que são encapsulados na PDU L6 e assim até a camada de rede com a NSAP, ou seja, o endereço da aplicação é constituído pela concatenação dos subendereços de cada camada (PSAP + SSAP + TSAP + NSAP).

Para o serviço Profile Service, com confirmação, ele envia um *request*, o usuário remoto recebe a mensagem *indication* e confirma (*response*), de forma que o usuário inicial receba a confirmação (*confirm*).

Para o *request* o usuário solicita um recurso e aguarda a resposta que será enviada através de texto formatado no padrão entendido pelos clientes (HTML), com o ciclo a seguir:

cliente (usuário *e-science*) => *request* => profile service => processa => *response* => cliente (usuário) => *confirm*

Feito então a definição do *layout* para a página do cliente *e-science*, o usuário primeiro terá que solicitar permissão, ao usar o serviço pela primeira vez, e adquirir seu acesso informando um perfil de acesso, ou seja, se é aluno, professor, pesquisador, etc., ele desta

forma vai realizar seu cadastro como usuário da infraestrutura e, feito isso, receberá pelo seu *e-mail* um *login* e senha. Ele pode escolher o tipo de acesso que desejar, ou mesmo criando novos, ele poderá criar seus projetos, realizar agendamentos dos arquivos que vai precisar transferir, como vemos nas Figura 37 e Figura 38, detalhadas abaixo.



Figura 37. Web service

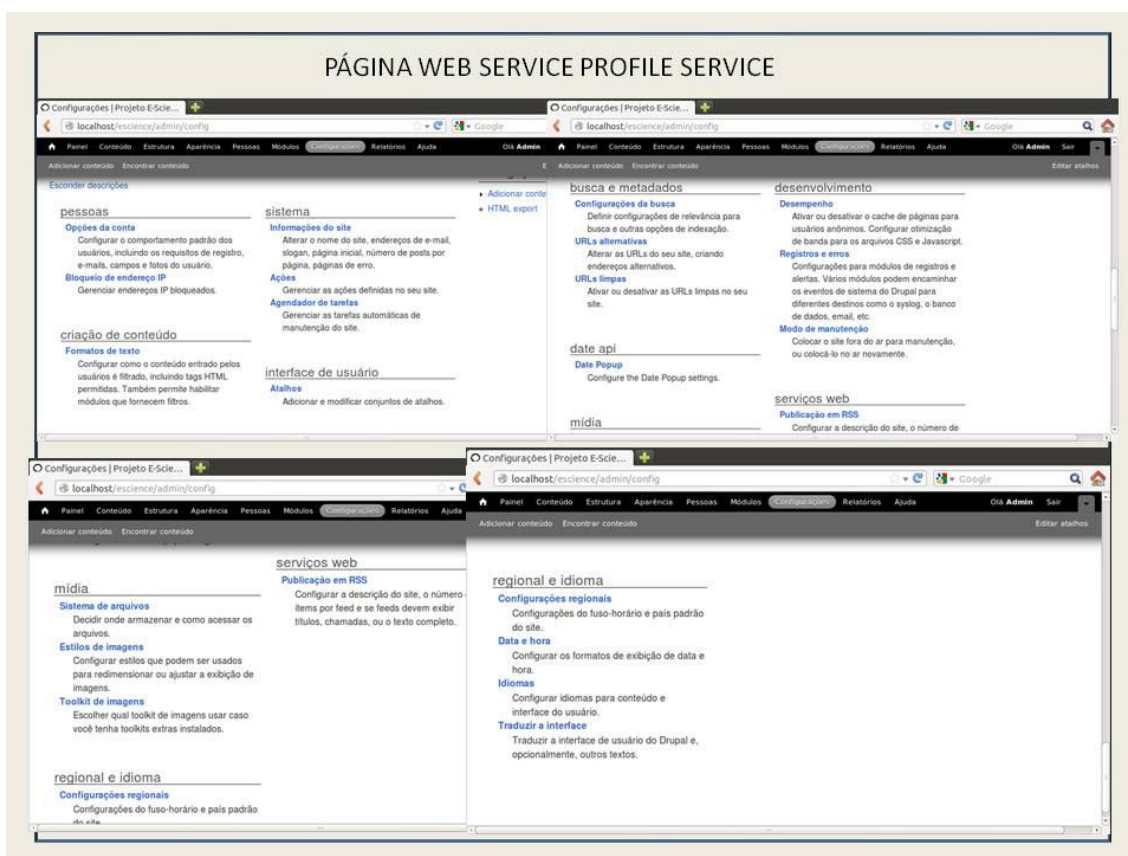


Figura 38. Visão geral

As primitivas trocadas são para estabelecer as conexões. A camada de enlace de dados encapsula o pacote em um quadro, adiciona um cabeçalho que inclui os endereços de origem, destino e FCS, para determinar se o quadro recebido é válido ou foi corrompido. Então este quadro é transmitido para a camada física que converte os dados em propriedades com base no tipo de cabo.

A camada de rede encapsula em um pacote e adiciona um cabeçalho que inclui uma camada lógica de endereços e o protocolo da camada superior. A camada de transporte encapsula os dados a partir das camadas superiores em um segmento de cabeçalho que inclui a porta de destino para receber a conexão, em seguida devolve esta informação para as camadas mais baixas.

O nível de rede deve prover os meios funcionais e processos para transmissão orientada a conexão entre as entidades do nível de transporte, que devera oferecer quanto ao roteamento associado ao estabelecimento e operação da conexão de rede, é função se tornar transparente de tal forma que as conexões do enlace sejam usadas para implementar as

conexões, executando procedimentos para “encobrir” as diferentes tecnologias de transmissão.

A confiabilidade dos pacotes inter-redes está diretamente ligada ao protocolo dentro da estrutura hierárquica. Nos serviços orientados a conexão um caminho lógico é estabelecido entre a origem e o destino.

Por conter mecanismos de controle de erro de fluxo, garante a sequência dos pacotes ao destino.

Nos serviços de pré-alocação de recursos, garantindo um melhor controle de congestionamento nos nós, esse serviço apresenta alto grau de complexidade, tempo de processamento e nem sempre a banda disponível.

Em nível de encaminhamento de pacotes, uma conexão virtual é aquela em que todos os pacotes seguem a mesma rota determinada na abertura da conexão.

A função do protocolo de inter-redes [56] é mover dados de um ponto a outro da rede. A camada de transporte é responsável pela movimentação dos dados de maneira eficiente e confiável entre os usuários. Esta camada deve regular o fluxo de dados e garantir a confiabilidade. Este protocolo deve possuir interfaces API e baseada em formulários (WBUI), que permitem que o usuário crie, pesquise, liste e exclua suas reservas. Devem oferecer segurança por meio de SSL, que criptografa as mensagens trocadas entre cliente / servidor.

O serviço de transporte oferece meios de estabelecer, manter e liberar as conexões de transporte entre um par de usuários, através dos seus pontos de acesso, para estabelecer a conexão de transporte, em uma rede com serviço orientado a conexão.

A camada de transporte tenta atender da melhor forma possível aos requisitos definidos pelo usuário e providencia uma identificação para a conexão, que uma vez estabelecida, começa a fase de transferência dos dados.

Para o serviço Profile Service o modo de encapsulamento será de acordo com a tecnologia da nuvem de transporte e será considerado um circuito virtual, se a rede for, por exemplo, ASON o circuito será um SNC; se a rede for SDH / OTN, será um VC ou um PWE3 para MPLS. Neste item é detalhado o circuito como um EVC para a rede *ethernet*.

O modo de encapsulamento Mac-in-Mac (802.1ah), quando os dados do usuário atinge o primeiro ponto da rede do provedor, este nó encapsula o MAC e adiciona um *tag* de *backbone*, chamado VLAN TAG ou B-TAG. Desta maneira os demais nós da rede passam a

processar apenas uma camada de endereçamento, sem precisarem aprender um grande número de endereços MAC.

O protocolo ainda trabalha com a proteção 1:1 para os caminhos pré-estabelecidos.

O ponto de entrada na rede do provedor transforma o *frame* do endereço c-MAC para o *frame* Mac-in-Mac, de acordo com o identificador do usuário. Este nó só processa o endereço b-MAC, ao chegar no destino o nó retira o endereço b-MAC e mantém o endereço c-MAC, conforme 2.4.1.

Com esta integração de tecnologias, ocorre uma proteção das informações dos endereços MAC dos clientes, além de redução de processamento no núcleo da rede, que não vai precisar aprender todos os endereços e resolve um problema de escalabilidade das redes *ethernets*. É então completado o processo de encaminhamento dos circuitos na rede.

Em relação a evento de falha, o serviço Profile Service prevê proteção no acesso, através da tecnologia *trunk*, que cria múltiplas interfaces físicas em uma interface lógica, e ainda o equipamento do cliente pode estar conectado a dois equipamentos do provedor da rede, em vez de estabelecer o tronco somente lógico.

Esta solução garante a confiabilidade do *link* e também no nível da rede, porque estará ligado a dois sistemas.

4.3.6. PRIMITIVAS DO SERVIÇO

A requisição para a utilização do serviço é feita pelo próprio cliente ou pela sua aplicação, através de um *front end* ou de uma API respectivamente, uma vez que já exista uma conexão criada. Esta conexão inicial é necessária, não apenas do ponto de vista da conectividade física em si, mas também para escoar a sinalização entre o usuário e a plataforma DCN.

Então, podemos definir que as primitivas (operações) do serviço são as entidades de transporte comunicando-se com seus usuários através de mensagens trocadas entre um ou mais pontos de acesso, que são definidas de acordo com o tipo de serviço e serão transportadas através das mensagens dos protocolos entre as interfaces. Para este serviço é especificado um conjunto de operações e as primitivas permite ao serviço, realizar ações ou reportar uma ação sobre uma entidade par [22].

São elas:

- *Request* - uma entidade solicita ao serviço um determinado trabalho.
- *Indication* - uma entidade deve receber uma informação sobre um evento.
- *Confirmation* - uma resposta a um pedido retorna à entidade.

A primeira primitiva do serviço indica a importância da informação a ser transmitida em relação a algum descarte da informação na rede e possui valores de prioridades. A segunda primitiva indica a uma entidade da camada superior a existência de congestionamento na rede e a terceira primitiva envia para a camada física de um elemento remoto, através da conexão física, a mensagem que deve ser transportada sem erros.

A mensagem de *request* inicia no cliente (da UNI-C → UNI-N) para a rede e vice-versa. É o pedido de inicialização da primeira conexão, estabelecendo a chamada. A mensagem seguinte é a *indication*, (da UNI-N → UNI-C) onde a sinalização indica qual a configuração para iniciar a chamada. Na fase de *confirmation* (da UNI-C → UNI-N), a sinalização confirma a chamada da primeira conexão e estabelece a chamada.

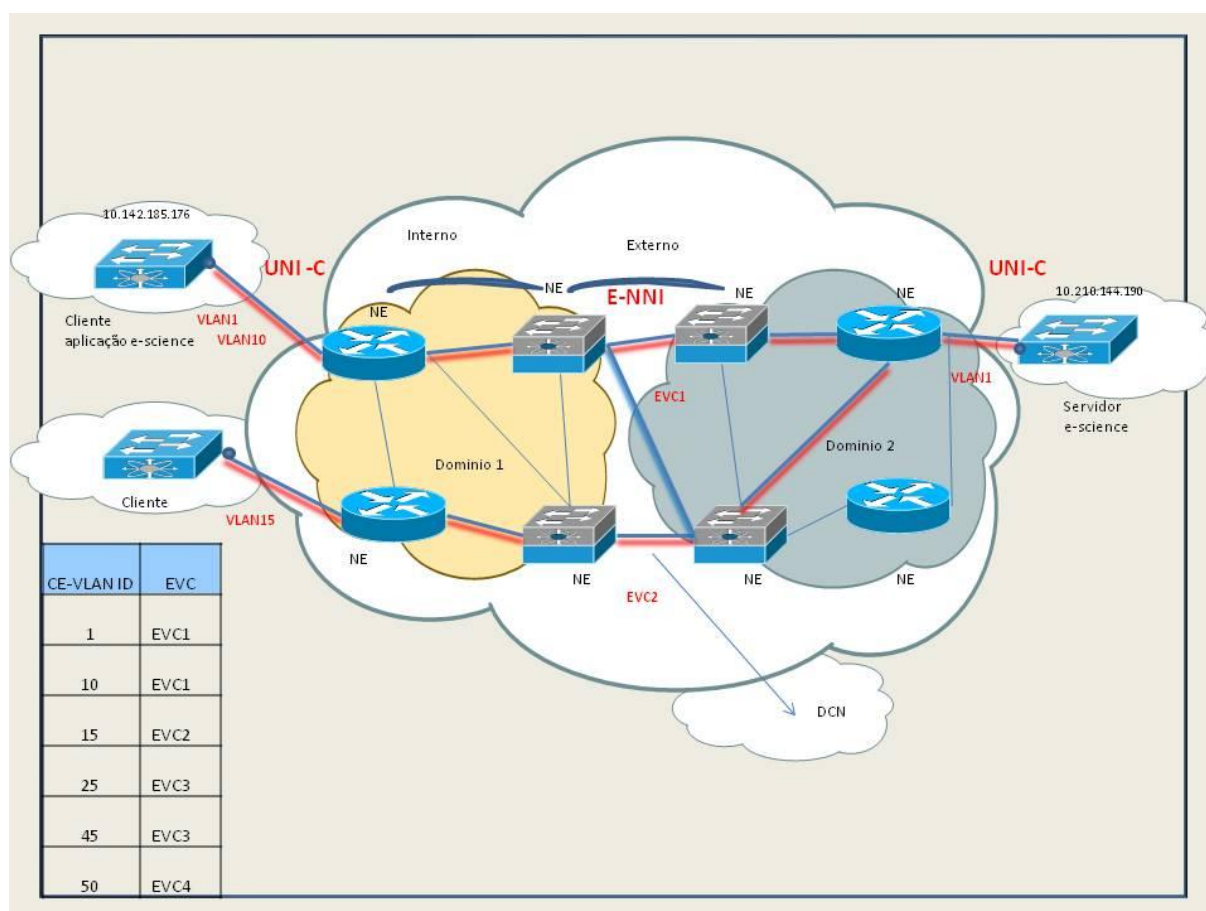


Figura 39. Mapeamento VLAN

Como podemos observar na Figura 39, para cada EVC na rede, há uma lista de UNI com identificador para cada circuito na EVC, desta forma, ao mapear o *frame* com o CE-VLAN ID, será transportado de acordo com as propriedades e atributos do EVC correspondente. Todo o serviço é baseado em VLANs com 2 níveis, a VLAN do usuário e a VLAN do *backbone*, chamadas cVLAN e bVLAN, onde o trio porta + cVLAN + dVLAN define o serviço que está sendo oferecido.

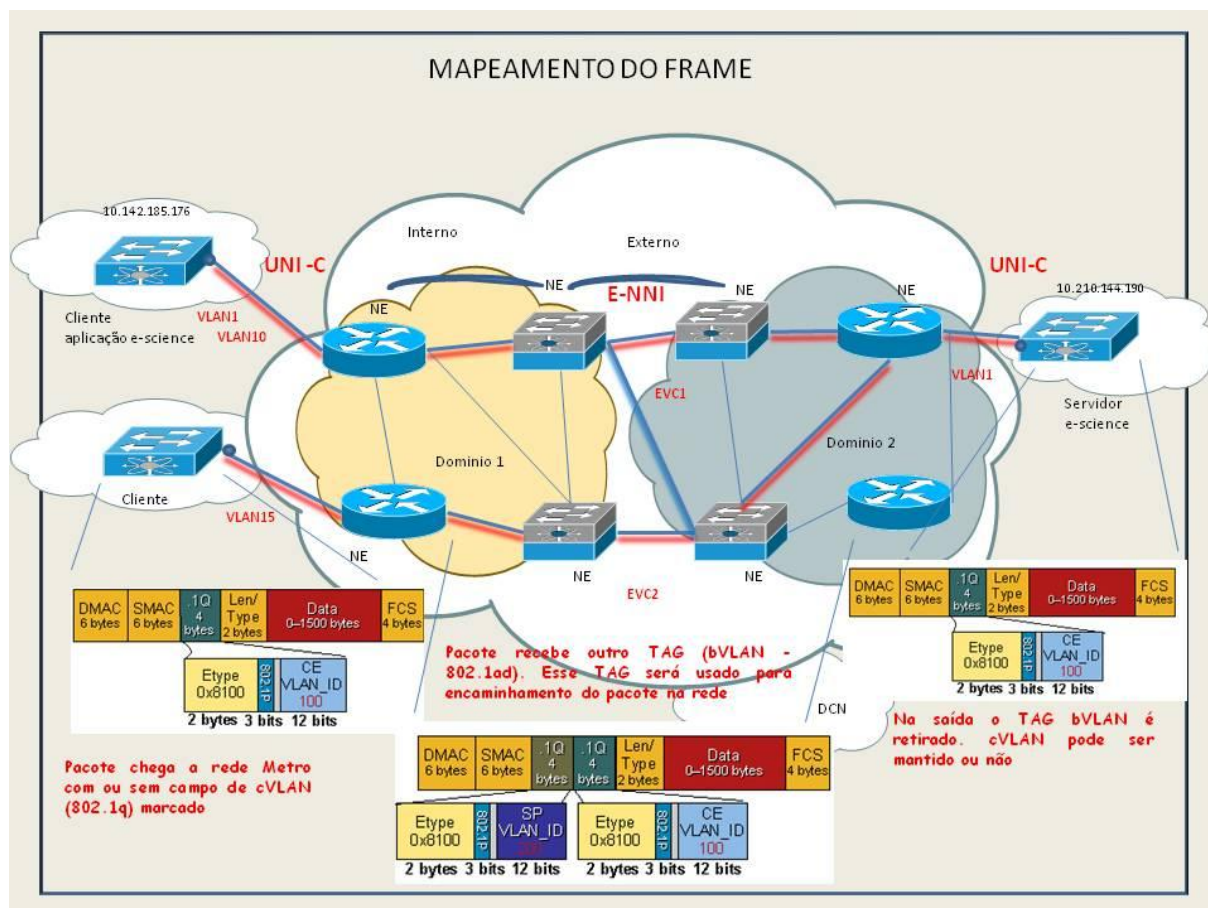


Figura 40. Mapeamento Frame

Quando o pacote do cliente chega na entrada da rede, ele é marcado no campo cVLAN e este, ao entrar no primeiro ponto da rede do provedor, recebe outro *tag*, agora para identificar o *backbone*, e é usado para o encaminhamento do pacote até a sua retirada, para ser então entregue ao equipamento do cliente da outra ponta. Um cliente pode demandar mais de uma bVLAN. Este mapeamento é realizado na camada 2 e está apresentado na Figura 40. A interface UNI permite que os usuários sinalizem para os serviços de transporte e possuem como características, o suporte a cliente *ethernet*, provendo serviços EPL e EVPL, modificação dinâmica de banda e conexão a serviços SDH e OTN. Ela especifica, ainda, o

protocolo que permite ao cliente a requisição e estabelecimento através da rede do provedor, mas sem ter visibilidade da rede da Operadora.

As mensagens da interface são:

- Criação de conexão – permite uma conexão entre pares de ponto de acesso;
- Apagar conexão – permite a supressão da ligação entre pares de ponto de acesso;
- Consulta *status* – permite ação consulta aos estados dos parâmetros da conexão;
- Modificação da conexão – permite alteração de parâmetro através de modificação adicionando ou removendo uma conexão existente.

A interface permite o transporte dos serviços *ethernet* através da conexão UNI-C a UNI-N, com os seguintes serviços definidos conforme a Figura 41.

RESUMO DOS SERVIÇOS UNI		
Tipo Serviço	Porta	VLAN
E-LINE	EPL	EVPL
E-LAN	EP-LAN	EVP-LAN
E-TREE	EP-TREE	EVP-TREE

Fonte: MEF 6

Figura 41. Serviços da UNI

Para os serviços agendados, através de um contrato com a Operadora, o cliente terá permissão para requisitar pedido de conexão em um determinado período (dia e hora), neste

caso, o agente da UNI-N precisa verificar se a conexão solicitada conseguiu, no agendamento realizado, receber o conteúdo contratado (Figura 42). Essas informações para a política de decisão estão contidas nas mensagens de sinalização UNI. Na interface do cliente, as mensagens de sinalização podem ser transportadas usando OAM, descrito no padrão IEEE802.ah.

Administração de Usuários e-science

Nome: eunice_aguiar
Senha: @e-science

Serviços Oferecidos
Consultas
Banco Dados
Solicitação de serviço
Troca de senha
Agendamento
Alterar agendamento

(Selecionar o serviço)

Figura 42. Seleciona o Serviço

Baseados no desenho apresentado para a arquitetura de aplicação, chegamos ao desenho de 3 modalidades (*Gold, Silver L2, Best Effort*) de oferta para os novos serviços, que o cliente seleciona conforme Figura 42.

4.3.7. MODALIDADES DO SERVIÇO

Este item apresenta as modalidades do serviço *ethernet*, conforme descrito abaixo na Figura 43.

Em um EVC (*Ethernet Virtual Connection*) ponto a ponto são necessários 2 pontos da interface UNI para serem associados, pois o *frame ethernet* ao chegar no nó *ingress* da rede será mapeado para o EVC na UNI *egress* e não em outro EVC. Da mesma forma deve ocorrer com um EVC multiponto, que deve ser associado a várias UNIs, ou seja, uma ou mais UNIs podem ser designadas como raiz e folhas associadas, mas somente uma delas como raiz de forma que o *frame* do serviço deve ser mapeado com estas designações.

Baseados no desenho apresentado para a nova arquitetura de aplicação e serviços, chegamos ao desenho de 3 modalidades (*Gold, Silver L2, Best Effort*) de oferta para o serviço, conforme a seguir.

REQUISITOS DOS SERVIÇOS			
CARACT. TÉCNICAS	GOLD	SILVER L2	BEST EFFORT
INTERFACE	100 Base FX, 1000 Base SX	100 Base FX, 1000 Base SX	100 Base FX, 1000 Base SX
BANDA GARANTIDA	30%	10%	0%
TEMPO RESTAURAÇÃO	< 50 ms	MODO STANDBY	SEM BACKUP
LATÊNCIA (ms)	5 ms	15 ms	30 ms
PERDA PACOTES (%)	0,001%	0,1%	0,5%
PRIORIDADE NA REDE	SEM PREEMPÇÃO	COM PREEMPÇÃO	SEM PRIORIZAÇÃO
DISPONIBILIDADE (%)	99,90%	99,90%	99,80%

Figura 43. Requisitos dos Serviços

Os diferentes tipos de serviços *ethernet* possuem conjunto de atributos que definem as características do serviço e estes atributos possuem parâmetros associados.

Para o serviço Profile Service, a Figura 43 apresenta as características técnicas que são os atributos relacionados com as 3 modalidades: de rede (interface e banda), de proteção

(restauração de disponibilidade) e QoS (latência, perda pacotes e prioridade CoS ID). O perfil de banda é caracterizado pelo quadro do serviço, com o propósito de policiamento da taxa de dados, que poderá ter 2 perfis, na entrada ou na saída de uma UNI.

Os parâmetros vão estar associados aos tipos de modalidades (Gold, Silver ou Best Effort) sendo que o mapeamento pela prioridade a ser usada na rede, pode ser feita pelo EVC e pelo *tag* 802.1Q.

O EVC é uma instância de associação de duas ou mais UNIs e os quadros só podem ser trocados entre UNIs associadas, com seus parâmetros de CIR e EIR associados, para os tráfegos com alta exigência de desempenho e os que não possuem estas mesmas exigências [51].

Cada UNI terá um perfil aplicado a todos os quadros de ingresso da UNI e cada EVC terá seu perfil de banda de ingresso definido pelo CoS ID, que pode variar de 0 a 7 dependendo do tipo de tráfego.

O desempenho da rede na entrega do serviço é especificado para os quadros dentro de um EVC e de acordo com o CoS. A disponibilidade é a percentagem do tempo durante o qual a taxa de perda fica dentro dos parâmetros definidos (Figura 33) e cada par de UNIs em uma EVC pode ter uma métrica diferente de desempenho.

Os serviços E-LAN permitem a definição de serviços VPNs e transparentes LAN, que vão do mais simples sem garantias de QoS entre as UNIs até os mais complexos, com garantias de desempenho para cada instância e que estão definidos abaixo:

Gold - Tempo real e plano de controle: Classe com a maior prioridade, uma vez que trafega dados essenciais para a operação da rede e dados de aplicações sensíveis a atraso e *jitter*. Especialmente em casos de redução da banda disponível devido a falhas, este tráfego não deve ser impactado, justamente para permitir a restauração do serviço e a recuperação da falha. Não sofre preempção e deve ter limiares de perda e retardo muito restritos, assim como a utilização de circuitos *backup* ativo, para rápida convergência, e outros mecanismos para restauração e recuperação automática. Utilizada para aplicação de videoconferência, por exemplo.

Silver L2 - Dados críticos em modo de transporte: Permite emular uma rede de transporte, provendo circuitos transparentes ao tráfego do usuário. Por este motivo não há diferenciação das aplicações que usam esta classe, a mesma política de QoS se aplica a todos

os pacotes, sem a possibilidade de diferenciação. Utilizada para aplicações e distribuição de conteúdo e transferência de dados.

Best Effort: Serviço *Best Effort* tradicional, sem a garantia de QoS ou de restauração/recuperação. Para aplicação de acesso à *Internet*.

4.4. A REDE

4.4.1. SUPORTE AO SERVIÇO DCN

Neste item vamos apresentar como os serviços terão os seus requisitos atendidos do ponto de vista do *backbone* do provedor da rede, que foi especificado ou tratado por Monteiro (referencia)[21].

Como proposta para solução desse problema, foram analisadas duas (2) alternativas, que utilizam o QoS para garantir que a DCN e as demais aplicações sejam separadas e priorizadas, de forma a garantir que o agendamento que o usuário solicitar, seja conforme sua necessidade.

A proposta é de alocação dos tráfegos em quatro (4) filas de QoS que vão distinguir e segregar por tipos de serviços, em uma fila exclusiva, da seguinte forma, a primeira classe destina-se ao tratamento de fluxos de tráfegos interativos e em tempo real, como voz e vídeo, para este tráfego o tratamento dado aos pacotes deverá ser o de menor tempo. A segunda classe de fluxos será diferenciada pela alta importância, associados a uma fila de alta prioridade, a terceira classe, corresponde aos fluxos que não precisam de nenhum tipo de marcação e a quarta fila seria exclusiva ao transporte dos fluxos de dados gerados pelo serviço DCN.

Para esta quarta fila, foi definida uma proposta de marcação dos pacotes e enfileiramento de QoS, diferenciando os fluxos de dados dos serviços, na qual eles terão a marcação no campo TOS no agrupamento AF1X, onde X pode ser igual a 1, 2 ou 3, dependendo da preferência de descarte, ou seja, quanto menor o valor numérico, maior a prioridade, de acordo com a Figura 44.

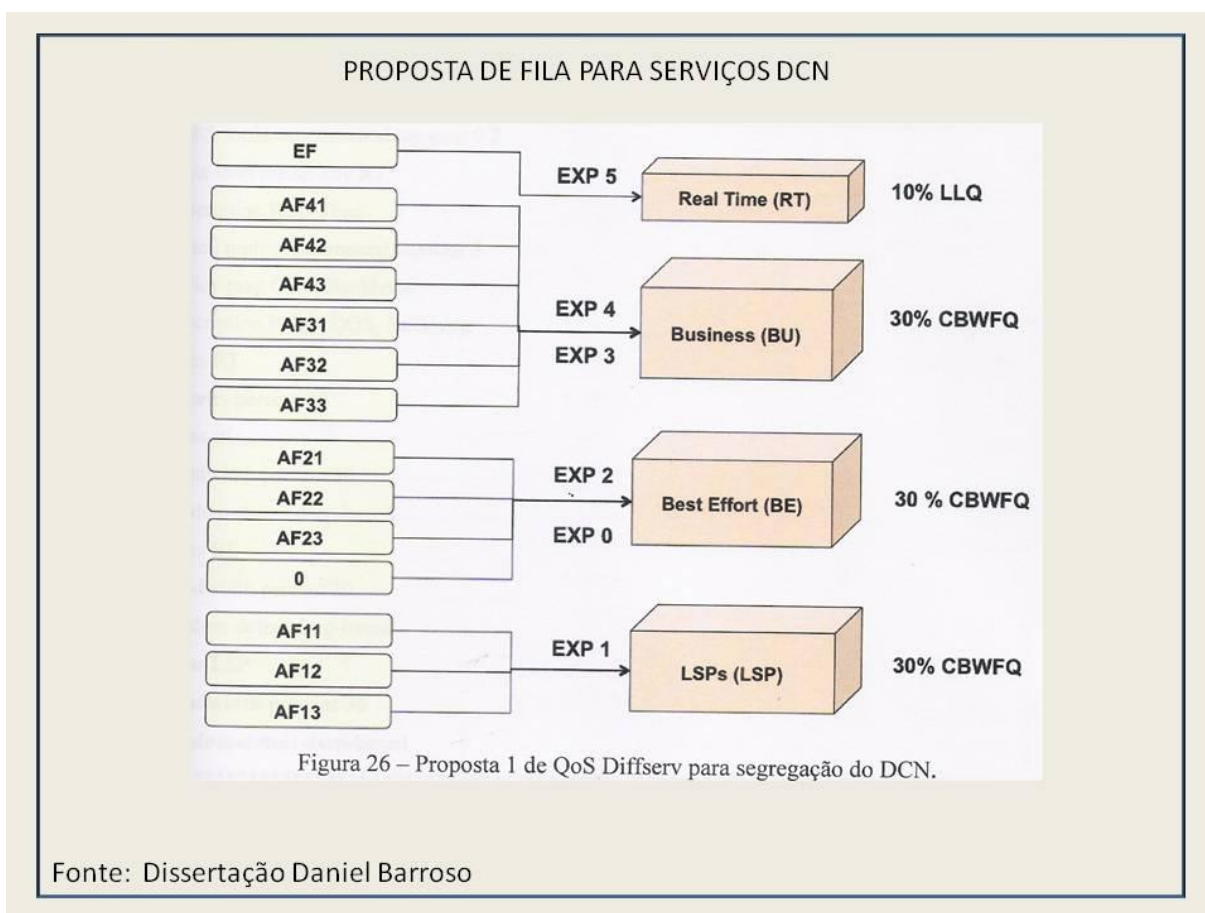


Figura 44. QoS para a DCN

A segunda alternativa é que poderia haver cinco (5) filas de QoS, o que nem sempre é possível pelo tipo do equipamento instalado na rede. Caso seja possível criar 5 filas, então seriam segregados os tráfegos de tempo real. Esta fila tem 10% de banda reservada.

Nesta dissertação são integradas as etapas da montagem do circuito na rede de forma fim a fim e de forma a mapear os serviços formatados no item 4.3.7. Neste trabalho, podemos verificar que o serviço *Gold* tem o tratamento do TIPO 1, os serviços *Silver* do TIPO 2 e o *Best Effort* do TIPO 3, conforme a tabela na Figura 45.

MODALIDADES DOS SERVIÇOS									
MODALIDADE DO SERVIÇO	CARACTERÍSTICA DO SUPORTE AO SERVIÇO DCN								
	LSP L2	LSP L3	QoS L3	QoS EXP	LSP Standby Ativo	LSP Standby Passivo	Fast Reroute	Prioridade do LSP	Exemplos de Aplicabilidade
TIPO 1	X			6	X		X	1	Emulação de circuitos transparentes ethernet, transporte de dados não TCP/IP de alta prioridade e baixo descarte, voz e vídeo
		X	X	1 e/ou 6	X		X	1	Transporte de diversos serviços IP de alta prioridade como voz, vídeo e automação industrial, simultaneamente
TIPO 2	X			1		X		3	Emulação de circuitos transparentes com média probabilidade de descarte, transporte de dados não TCP/IP não interativos ou de tempo real
		X	X	0 e/ou 1		X		3	Transporte de diversos serviços IP de média prioridade como aplicações de missão crítica não interativas ou de tempo real
TIPO 3	X			1				5	Emulação de circuitos transparentes com alta probabilidade de descarte, transporte de dados não TCP/IP de baixa prioridade
		X	X	0 e/ou 1				5	Transporte de diversos serviços IP de baixa prioridade que podem ser descartados preferencialmente

Quadro 6 - Modalidades de suporte ao serviço DCN e suas características.

Fonte: Dissertação Daniel Barroso

Figura 45. Modalidades

Os serviços Profile Service tipo *Gold* são mapeados com o *frame* cVLAN com o identificador CE VLAN ID e na entrada da rede receberá outro *tag* bVLAN do *backbone* do provedor é o SP VLAN ID do tipo 1 com o QoS AF11, desta forma fica a garantia para o serviço do cliente, que ele vai percorrer toda a rede com os requisitos garantidos pelo tratamento da rede, na priorização dada pela fila exclusiva da DCN.

Nas demais modalidades, *Silver* e *Best Effort*, é a mesma identificação de mapeamento dos *frames* na rede, sendo que para os serviços *Silver* ficam no Tipo 2 com o QoS AF12 e o *Best Effort* com o Tipo 3 e QoS AF13.

4.4.2. IMPACTOS NA REDE

Uma vez que a rede é existente e operacional, é preciso identificar os riscos para os demais serviços já prestados. O principal risco, que pode ser percebido facilmente, é a

possibilidade de vir a causar *starving* aos demais serviços, uma vez que a interface com o usuário, por ser *ethernet*, permite a sua utilização acima da velocidade nominal da banda reservada.

A proteção aos demais serviços já prestados deve ser feita por duas vertentes simultaneamente, uma contratual, prevendo restrições ou custos extras para o tráfego excedente, e outra técnica, configurando limitações na própria rede. Como já sinalizado deve ser feito um cuidadoso estudo de capacidade da rede para determinar quanto pode ser usado pelo serviço Profile Service sem trazer problemas aos outros serviços. Para isso é preciso levantar:

- Banda nominal nos vários enlaces da rede.
- Banda disponível, considerando dados atuais e históricos.
- Demanda atual e prevista para os serviços existentes.
- Demanda atual e prevista para o Profile Service.
- Capacidade dos equipamentos usados na rede e a sua possibilidade de expansão.

Por outro lado é importante definir o nível de serviço a ser ofertado em termos de disponibilidade e confiabilidade, considerando os tempos médios e máximos de reparo, tempo médio entre falhas e garantia de restauração do serviço. Devem ser avaliadas as garantias de QoS de cada modalidade: latência, perda de pacotes, banda, tratamento dado ao tráfego excedente , etc., de acordo com a sua prioridade. Finalmente é preciso incluir questões acerca do uso de *overbooking*, bem como limiares de utilização de enlaces e equipamentos tanto no *backbone* como no acesso. A partir daí poderá ser definida a melhor técnica para fazer a proteção.

Recomendamos o uso do *DiffServ*, para separar em diferentes filas do Profile Service e os demais serviços, subdivididos nas suas respectivas modalidades. O novo serviço e o legado estariam assim isolados obrigando-os a utilizar os recursos conforme a política definida pelo gerente da rede. O RSVP-TE pode então ser usado para fazer as reservas compartilhando os recursos alocados para o Profile Service entre seus usuários. As solicitações serão inteiramente reservadas, independente da sua utilização ou não, e esta informação será usada para determinar a disponibilidade de recursos para a sinalização de um próximo circuito. Ainda que haja banda disponível se houverem reservas consumindo completamente a banda destinada ao Profile Service, novos circuitos não serão sinalizados. Caso o novo circuito tiver

maior prioridade, serão seguidas as regras de preempção, redefinição ou mesmo remoção definidas para cada modalidade.

Além do uso dos recursos da rede, outro impacto importante que deve ser mensurado são os riscos de segurança, que o novo serviço pode trazer para a rede. Devem ser avaliadas as vulnerabilidades envolvendo diferentes fluxos de tráfego, diferentes usuários e a rede em si. Devem ser protegidos o conteúdo e a integridade dos fluxos de tráfego, das redes virtuais sobrepostas à rede do provedor, da própria infraestrutura de rede e de seus sistemas, como bilhetagem, gerência, provisionamento e etc. Para isso é essencial garantir o isolamento entre as VLANs e usar protocolos de controle de acesso, para fazer autenticação, autorização e *accounting* dos usuários e suas ações, seja diretamente seja via plataforma de DCN, com a solicitação de circuitos.

Como resultado dos testes realizados por Monteiro [21] foi possível concluir que as propostas apresentadas e emuladas em equipamentos Cisco, possuem aplicabilidade em outras redes de outros fabricantes.

4.5. CENÁRIOS

Para o estabelecimento da conexão são necessárias informações corretas de forma que a sua construção seja atingida, e grande parte destas informações refere-se à identificação do usuário, de forma que o plano de gerenciamento possa comutar na direção correta.

Este plano de sinalização se inicia no usuário quando requisita o serviço e informa os parâmetros necessários para a sua aplicação, conforme a Figura 43. O gerenciador do domínio avalia os recursos e se os parâmetros são possíveis realizando um *check* local, devolvendo os parâmetros do circuito a ser criado, como o *delay* e o número da VLAN de volta ao usuário. Esta construção será realizada por todos os domínios existentes, sendo que os dados retornarão a partir do último domínio.

Vamos então resumir em três cenários, o atendimento ao cliente que demande um dos serviços *e-science* desenvolvidos no item 4.3.7, considerando o cliente como um domínio particular interagindo com o domínio da rede do provedor; em um segundo cenário o cliente possuindo seu próprio domínio DCN; e no terceiro cenário, o cliente participando do domínio DCN do provedor, como apresentaremos a seguir.

Para estabelecer os cenários, foi realizado estudo da aplicação *e-science*, do serviço dinâmico e da rede híbrida DCN, pois o conjunto de facilidades adquiridas com o uso das redes DCN os *frameworks* DRAGON e OSCARS, o plano de controle de alocação dinâmica e o uso de políticas de QoS, implementadas diretamente no *backbone*, motivaram o desenvolvimento deste trabalho.

A partir da abrangência do domínio de controle da rede DCN, foram criados cenários para análise. Num primeiro cenário, apenas o domínio DCN está presente na rede do Provedor, com acesso da rede do cliente de forma estático.

No segundo cenário temos vários domínios DCN implementados, tanto na rede do Provedor como na rede do cliente. Por fim, no terceiro cenário, com um único domínio DCN que engloba a rede do cliente e como único domínio de controle.

4.5.1. CENÁRIO 1

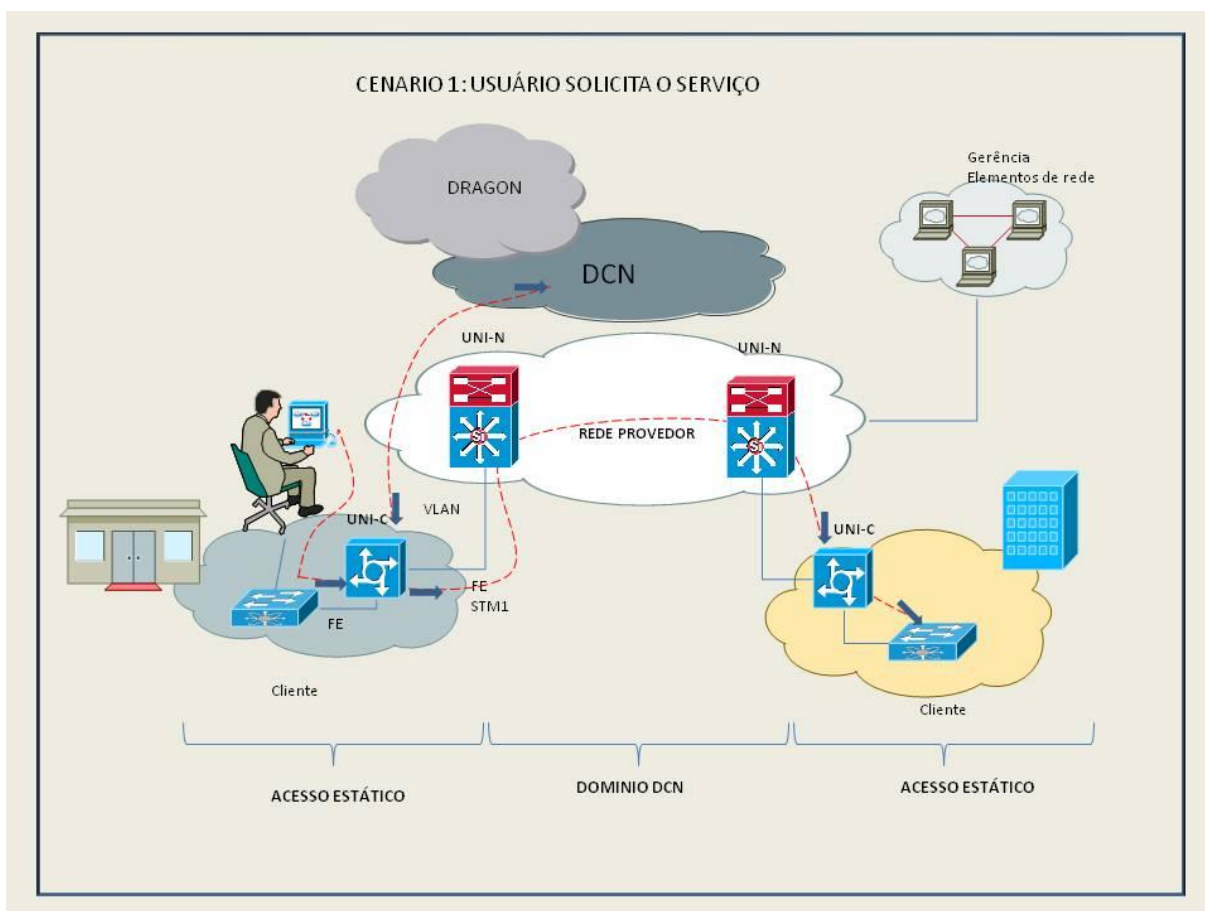


Figura 46. Cenário 1 - 1 domínio DCN

Para o cenário 1 (Figura 46) o acesso do cliente a rede do Operador pode ser através de equipamentos *ethernet* e a interface *fast ethernet*, nas velocidades definidas no item 4.3.3 e na rede do provedor podemos ter as tecnologias SDH, DWDM ou GMPLS, que farão a agregação na sua interface UNI-C.

Neste cenário observamos que o cliente pode acessar o site do provedor através da sua interface de cliente (UNI-C), que requisita o *setup* do serviço, para a interface de rede (UNI-N). Esta envia a solicitação para um servidor da rede e este para o plano de controle da rede, que avalia o volume das informações e o ajustar as necessidades do cliente com a disponibilidade dos recursos de rede, podendo alocar mais ou menos banda.

O usuário via aplicação *e-science*, seleciona os campos com os seus requisitos e então terá criado a sua reserva. Será necessária a validação do cliente, através de acesso *login* e senha, antes que o cliente receba um número de protocolo com o atendimento da sua solicitação. Para o cliente a rede será um circuito virtual (EVC), e o requisito de banda será solicitado através da gerência do *site* do cliente.

Neste cenário temos um domínio da rede DCN, que sé a rede do provedor, para o cliente o seu acesso será estático, os *softwares* DRAGON e OSCARS estarão implementados apenas na rede DCN. O protocolo no domínio da rede do provedor ao receber a solicitação do cliente, analisa e verifica se a rede atende aos requisitos do serviço solicitado e envia ao destino. Este resolve se o pedido será ou não aceito e envia mensagem de confirmação. Esta sinalização neste cenário envolverá o DRAGON no domínio da rede do provedor e o OSCARS entre o cliente e a rede, através de mensagens como: aceito, criado, pendente, ativado, modificado, cancelado, etc..

O protocolo encapsula em mensagens, após avaliar os parâmetros da solicitação, processa o pedido e retorna mensagem *<forwardresponse>* [56] e finaliza, conforme demonstrado na Figura 47 abaixo.

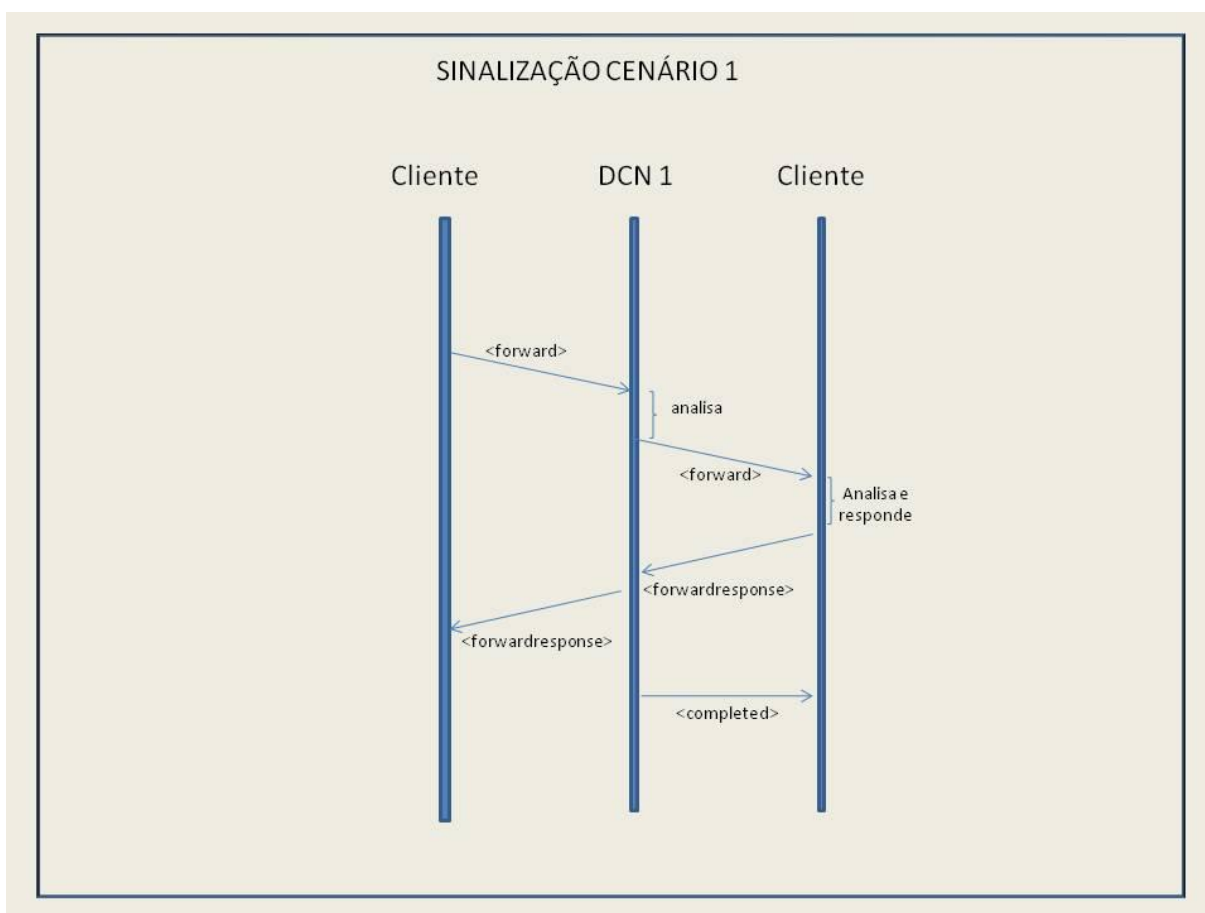


Figura 47. Sinalização cenário 1

Num primeiro momento, o cliente solicita via *website* o serviço Profile Service. O operador de rede estabelece o circuito na rede. Neste cenário não usaremos o OSCARS além do domínio do cliente. O Dragon vai aprovisionar o circuito intradomínio e garantir a autenticação.

A desvantagem deste cenário é a limitação para o usuário da utilização da comutação dinâmica, uma vez que o plano de controle DCN não vai estar no ambiente do cliente. Não recomendamos para este cenário a modalidade Gold ou Tipo 1 com alta prioridade, uma vez que o circuito não será criado fim a fim.

4.5.2. CENÁRIO 2

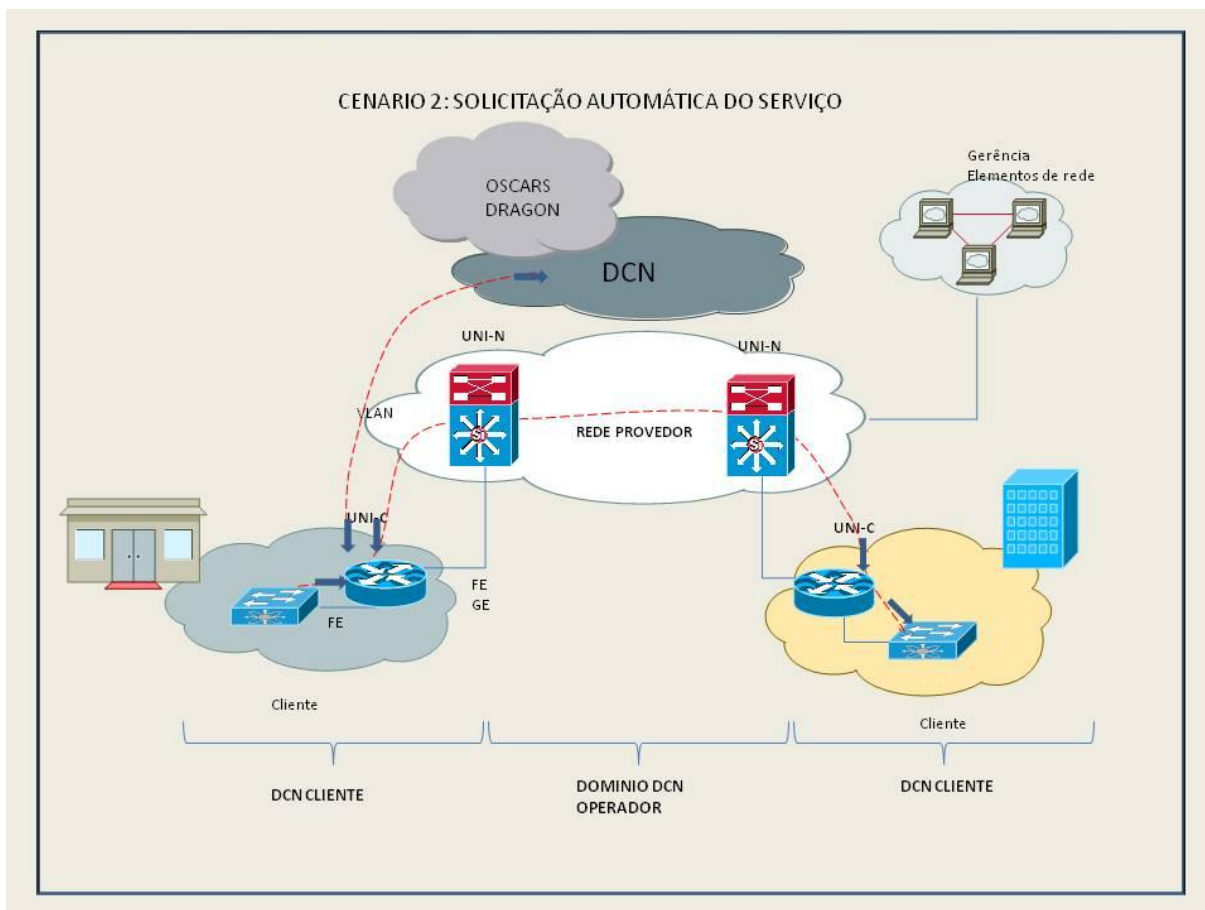


Figura 48. Cenário 2 - 3 domínios DCN

A mesma estrutura do protocolo interdomínios explicado no cenário será utilizado nos demais cenários, variando a abrangência do domínio de cada DCN.

Na Figura 48, o cliente acessa o *site* do provedor e, através da sua interface de cliente (UNI-C), requisita o *setup* do serviço, por meio do seu próprio domínio DCN, a interface de rede (UNI-N) envia a solicitação para um servidor da rede do domínio DCN do provedor e para o plano de controle, avalia o volume das informações e o ajusta para o serviço, podendo alocar mais banda ou reduzir a banda e envia para o cliente final que vai resolver se o pedido será ou não aceito e envia mensagem de confirmação ao domínio DCN seguinte a ele, e este analisa e modifica os parâmetros e devolve ao domínio anterior, neste cenário o domínio DCN do cliente de origem.

O usuário ao preencher os campos com os seus requisitos terá criado a sua reserva de forma automática. Será necessária a validação do cliente, através de acesso *login* e senha, antes que o cliente receba um protocolo com o atendimento da sua solicitação.

O cliente acessa a rede do operador através de equipamento de baixa ordem utilizando interfaces *fast ethernet* de forma automática através do protocolo interdomínio que encapsula o *frame ethernet* mapeando para a fila específica DCN com o QoS apropriado. Para o cliente, a rede será um circuito virtual, EVC.

Neste cenário temos vários domínios da rede DCN, que será a rede do provedor e a rede do cliente, os *softwares* DRAGON e OSCARS estarão implementados nos três domínios de rede, onde o processo será iniciado com o envio da solicitação pelo usuário para o IDC do primeiro domínio. Este encapsula em mensagem para o domínio seguinte, conforme a Figura 49.

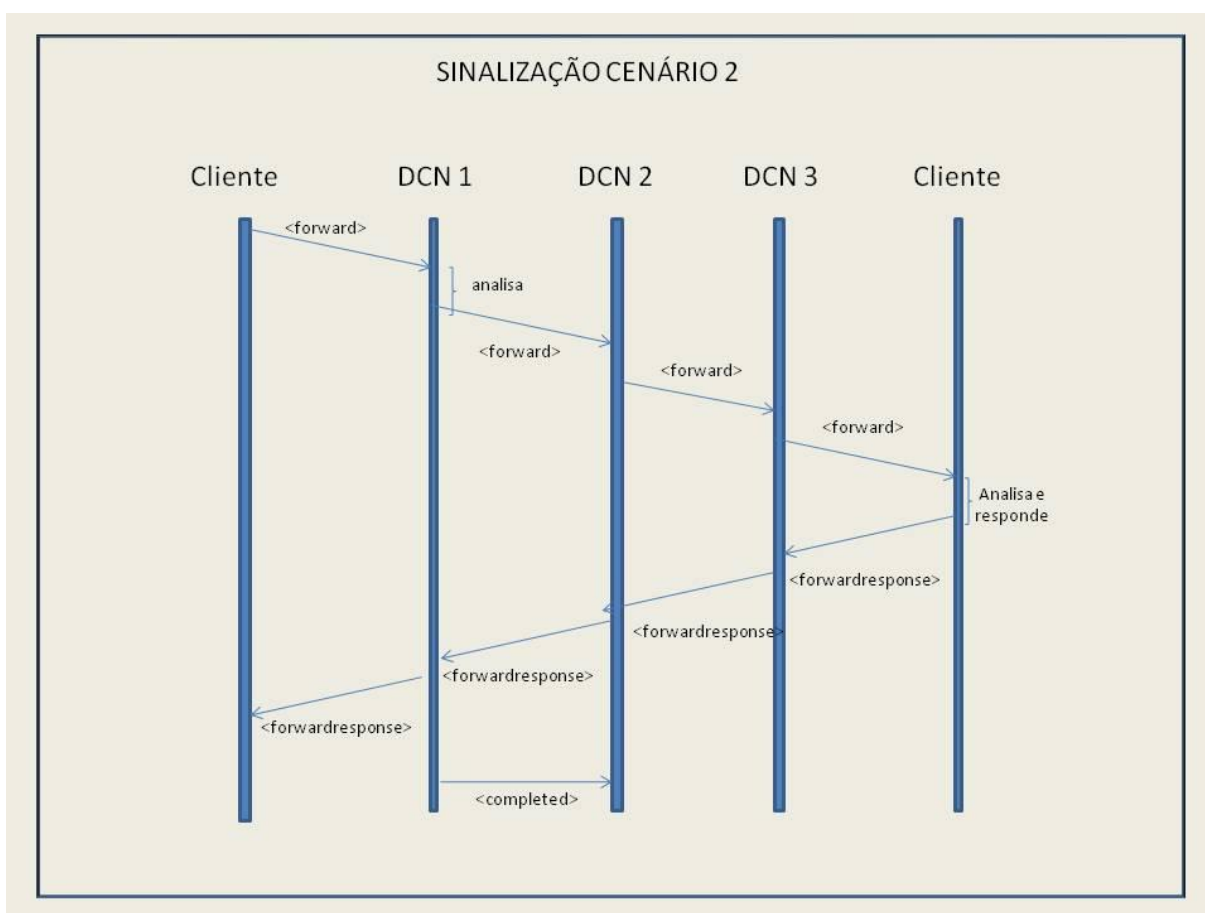


Figura 49. Sinalização cenário 2

Na Figura 49 podemos observar que o domínio 2 extrai os parâmetros da solicitação e torna a encapsular, enviando para o domínio seguinte, que vai processar o pedido e retornar mensagem de volta, que irá passar pelo domínio 2 chegando ao domínio 1 com a resposta ao pedido inicial do usuário.

A desvantagem deste cenário é o aumento da complexidade e o estabelecimento de circuitos intra e interdomínio, padronizados como os protocolos DC e IDC (item 3.2), mas com o uso do OSCARS e DRAGON, para aprovisionar os circuitos inter e intradomínios e garantir que os parâmetros de qualidade sejam fornecidos pela rede. Testes de validação do serviço Profile Service com OSCARS e validação da configuração do usuário mais profunda devem ser realizados antes de uma operação comercial. Neste trabalho realizamos testes preliminares de conectividade e prova de conceito com resultados favoráveis, indicando viabilidade técnica de tal cenário.

Para este cenário todas as modalidades podem ser ofertadas, pois o domínio DCN está presente fim a fim [56].

Neste cenário temos um único domínio de rede DCN, que é a rede do provedor até a rede do cliente, os *softwares* DRAGON e OSCARS estarão implementados no domínio único de rede (Figura 51).

A desvantagem deste cenário é que o cliente participa do mesmo domínio DCN do provedor, e então aparecem as questões de segurança, até que ponto da rede será feita a sua administração. Para o cenário 3 que expande a infraestrutura da rede até o cliente, temos as questões de acesso e segurança na rede, por estas razões, não recomendamos a sua implementação para nenhuma das 3 modalidades.

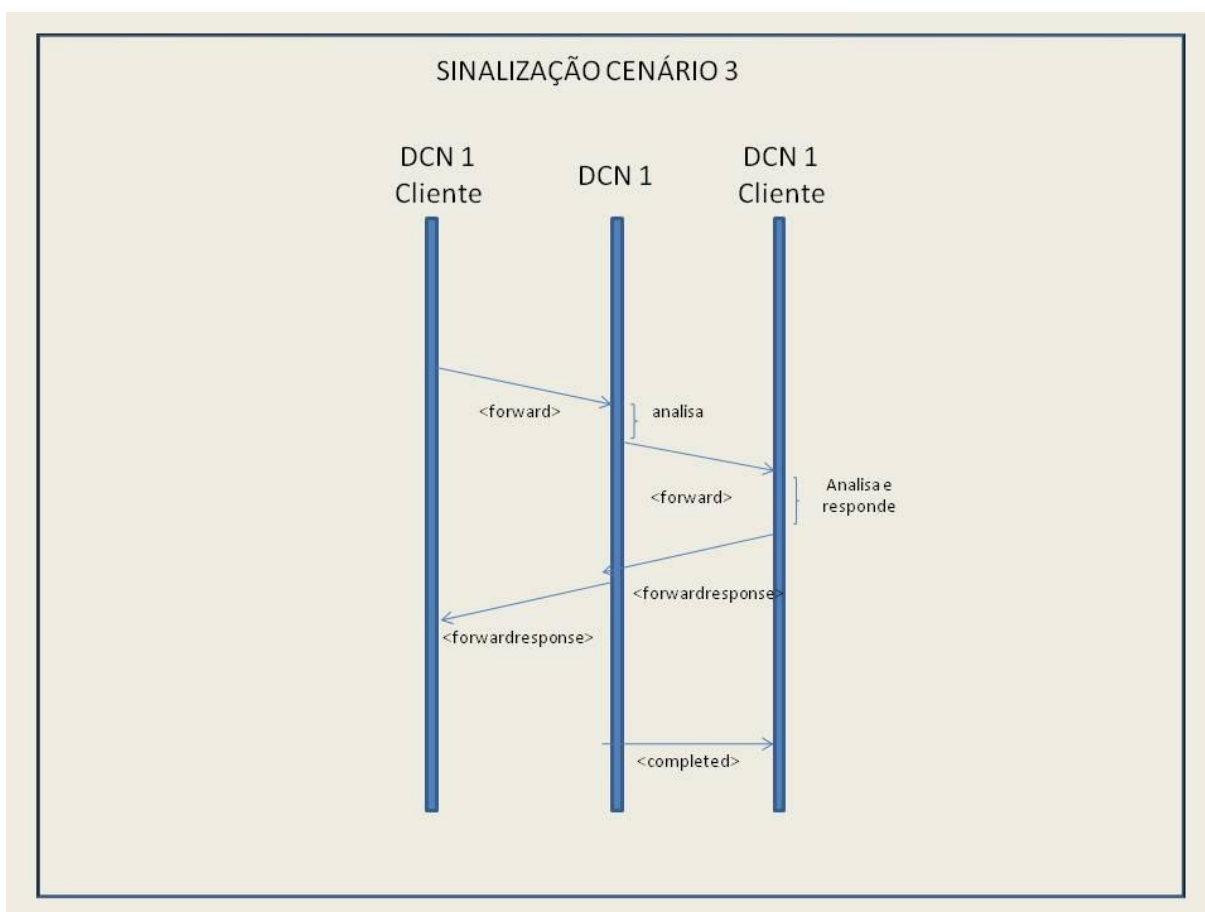


Figura 51. Sinalização cenário 3

Os modelos de serviços Profile Service já mencionados (item 4.3.7), podem ser implementados nos cenários descritos acima, porém, nesta dissertação recomendamos que seja adotado o cenário 2, com as modalidades Tipo 1, 2 e 3.

Para os três cenários descritos acima, poderão ser reavaliados os testes dos serviços, quando as novas placas dos equipamentos, possuírem a interface de cliente UNI 2.0 implementada, e então oferecer os serviços dinâmicos de forma automática e fim a fim. Os serviços serão mapeados com definição de QoS em uma fila específica nos roteadores do provedor, definida como EXP 1 CBWFQ, de forma a garantir que o tráfego da rede DCN, fique separado do tráfego da rede hoje em operação no provedor.

Nesse cenário, onde o dispositivo do usuário estará ligado ao domínio da rede de transporte, via interface do cliente, o protocolo RSVP vai definir as sessões na interface e serão associadas ao usuário final, através do endereçamento da fonte e do destino. O dispositivo do cliente pode ser atribuído ao seu próprio endereçamento, que não será referenciado na rede de transporte, permitindo a separação dos espaços de endereço.

Temos, então, um modelo de rede de transporte com base em multidomínios GMPLS, ligados via interfaces UNI, ao domínio da rede de transporte, que será composta por subdomínios. Serão criados 3 EVCs (AF11, AF12 e AF13) e a sessão I-NNI será definida pelo túnel estabelecido na rede específico para a rede DCN.

4.5.4. COMPARAÇÃO DOS CENÁRIOS

Cada um dos 3 cenários de rede DCN avaliados neste Capítulo apresenta suas métricas de desempenho para os tipos de serviço Gold, Silver e Best Effort. A medição do desempenho dos serviços apresentada na Figura 33 traz valores médios que serão transportados nas redes de transporte, sejam de circuitos ou de pacotes.

Na validação da boa transmissão dos serviços é necessário avaliar os atrasos dos pacotes, taxa de perda e disponibilidade, principalmente, além das questões de segurança nas redes DCN.

O serviço Gold, de alto valor agregado, fica melhor formatado nas redes do cenário 1 de acordo com a Figura 22 de novos modelos de rede, pois as redes de circuitos apresentam baixos valores de *delay*, *jitter* e alta disponibilidade. Os serviços Silver e Best Effort ficam melhor formatado nos cenários 2 e 3.

No entanto, as redes do cenário 1 com núcleo SDH / ASON possuem operação mais complexa, gerências distintas para cada trecho da rede, algoritmo de convergência complexo. Isto leva a maiores custos na operação da rede e monitoração do sistema.

O cenário 2 com núcleo MPLS / ETH possui controle de fluxo, mas difícil de realizar o sincronismo devido ao fato dos pacotes não seguirem o mesmo caminho entre a origem e o destino, necessitando de roteamento nível 3.

As novas tecnologias híbridas do cenário 3 possuem sincronismo, OAM simplificado e regras de engenharia de tráfego de camada 2 e com menores custos.

Então os serviços modelados ficam melhor formatado neste último cenário, em conformidade com as métricas de desempenho dos serviços.

4.6. APLICATIVO DESENVOLVIDO

A aplicação *e-science* desenvolvida para esta dissertação com o nome de Profile Service, é o *website* do cliente, ou seja, um conjunto de hipertextos acessíveis ao protocolo HTTP (*Hypertext Transfer Protocol*) e organizados a partir de um URL básico, onde fica hospedada a página principal e reside o diretório em servidor do conteúdo *e-science*.

Foi utilizada a ferramenta Drupal, um *framework* que gerencia conteúdo escrito em PHP (*Hypertext Preprocessor*) que independe do sistema operacional utilizado, mas que requer um servidor HTTP. O *software* Drupal implementa APIs com estrutura modular com licença GNU (*Linux*) [67].

4.6.1. ANÁLISE DOS SERVIÇOS UFF E RNP

Conforme mencionado anteriormente, o serviço SE-Cipó opera realizando chamadas da ferramenta OSCARS, onde os usuários através do *site* da RNP solicitam reservas de circuitos informando os pontos de origem e destino, a banda necessária e o intervalo de tempo que o circuito ficará ativo. Estas reservas são feitas através da ferramenta Meican com uso de mapas.

Cada requisição gera um pedido de autorização baseado nas políticas do domínio a que pertence o usuário. Para que se complete a etapa, todos os domínios envolvidos devem aceitar as condições da requisição.

A interação com a ferramenta OSCARS foi desenvolvida pela Esnet e através de convenio colaborativo, foi adequada ao Meican, chamado então de *middleware oscar_bridge*, desenvolvido em Java.

A RNP também vem trabalhando em uma interface. A arquitetura ficou estabelecida com o Meican interligado ao OSCARS apenas no *backbone* da RNP e cada POP com seu OSCARS que via protocolo IDC que comunicam em topologia estrela com o OSCARS principal, de forma hierárquica dentro da rede. A modelagem da aplicação da RNP ainda está em fase de estudo.

O serviço Profile Service desenvolvido neste trabalho usa o protocolo HTTP, através de um programa para o cliente com interface *web*. Quando o usuário abrir a URL do serviço, seu navegador abre uma conexão com o servidor na máquina principal do serviço *e-science*, que vai exibir as opções de banco de dados.

A ferramenta OSCARS usa módulo para configurar os equipamentos através do protocolo GMPLS, usando o SNMP para identificar comandos via Telnet e através do DRAGON nos sites dos clientes. Para realizar o processo de reserva do circuito, o OSCARS precisa configurar VLAN e conhecer a origem e o destino.

Então, o usuário primeiro terá que solicitar permissão, ao usar o serviço pela primeira vez, e adquirir seu acesso informando um perfil de acesso, ou seja, se é aluno, professor, pesquisador, etc., ele desta forma vai realizar seu cadastro como usuário da infraestrutura e feito isso, receberá pelo seu *e-mail* um *login* e senha e ele já estará podendo escolher o tipo de acesso que desejar, ou mesmo criando novos, ele pode criar seus projetos, realizar agendamentos dos arquivos que vai precisar.

Para o serviço o modo de encapsulamento será de acordo com a tecnologia da nuvem de transporte, e será considerado um circuito virtual.

As características técnicas que são os atributos relacionados com as 3 modalidades, de rede (interface e banda), de proteção (restauração de disponibilidade) e QoS(latência, perda pacotes e prioridade CoS ID). O perfil de banda é caracterizado pelo quadro do serviço, com o propósito de policiamento da taxa de dados, que poderá ter 2 perfis, na entrada ou na saída de uma UNI.

Os parâmetros vão estar associados aos tipos de modalidades (Gold, Silver ou Best Effort) sendo que o mapeamento pela prioridade a ser usada na rede, pode ser feita pelo EVC e pelo *tag* 802.1Q.

Avaliamos então, que em termos de arquitetura de rede, a RNP adotou o Cenário 2. Neste cenário temos vários domínios da rede DCN, que será a rede backbone da RNP e a rede

do cliente, o *software* OSCARS estará implementado nos três domínios de rede, onde o processo será iniciado com o envio da solicitação pelo usuário para o IDC do primeiro domínio.

Mesmo para os circuitos dinâmicos é importante prever limites de utilização, uma vez que os requisitos definidos na sua solicitação, implementarão modificações diretamente na estrutura do *backbone*, no encaminhamento dos pacotes e a criação ou desconexão de circuitos existentes. Deverá ser estudada uma forma de limitar a banda consumida pelos circuitos criados dinamicamente para que estes não concorram com os outros serviços providos pela rede da RNP. Uma das possíveis técnicas propostas é o uso do protocolo RSVP-TE com uma limitação percentual da banda, por interface dos roteadores, para a criação destes circuitos. Por exemplo, para um *link* entre roteadores de núcleo da rede MPLS a criação de circuitos pode ser limitada em até 30% da banda. Apenas poderão ser criados, circuitos até este limite, a partir de 30% de ocupação o próximo circuito não conseguirá ser sinalizado. Por conseguinte, 70% da banda deste *link* serão protegidas e liberadas para tráfegos de outros serviços da RNP.

O uso do protocolo RSVP-TE permite um grau de controle maior por parte da equipe de administração do *backbone*, na qual serviços importantes podem ser definidos para ocuparem circuitos de mais alta prioridade e terão preferência dentro da rede. Podem ser estudadas ainda limitações envolvendo quantidade, duração, banda máxima e/ou mínima, QoS máximo e/ou mínimo e uma combinação dos itens anteriores, por exemplo circuitos com banda maior do que xx Mbps por no máximo yy horas. Novamente a definição destas e outras limitações dependem de um estudo de capacidade da rede, ou outra funcionalidade (*feature*) que o equipamento do *backbone* tenha, e possa vir a ser aproveitado.

Para finalizar a questão de utilização do serviço Profile Service, devemos considerar a possibilidade de criação pelo cliente de uma rede virtual própria, com uma topologia virtual sobreposta à rede da RNP. Nestes casos pode ser previsto um nível de administração da rede virtual, subordinado à administração e operação da RNP, permitindo ao cliente adicionar novos pontos e circuitos. Todas as atividades realizadas pelo cliente devem ser feitas através de um controle de acesso. Protocolos de controle de acesso aos equipamentos limitam os comandos que podem ou não serem executados e devem ser implementados nos elementos de rede. Desta forma, os *scripts* de configuração gerados pelos cálculos computacionais de

criação de circuitos dinâmicos, por exemplo, serão limitados a executar o estritamente permitido nos elementos de rede.

4.7. RESUMO DO CAPÍTULO

O Profile Service tem como principal objetivo implantar um serviço com reserva da banda pelo cliente. É um sistema que permite estas reservas automáticas para criação de circuitos fim a fim, através de vários domínios.

É um serviço de comunicação de dados que utiliza a arquitetura das redes híbridas e comutação dinâmica para prover circuitos ponto a ponto de altas capacidades, com a solicitação a partir do cliente. Com proposta de forma organizada este serviço traz um plano de planejamento e de implementação, descrição das primitivas do serviço e políticas, que organizam as lógicas do serviço, topologia e QoS.

Foi planejado com 3 modalidades, que possuem um conjunto de atributos que definem as suas características (Gold, Silver L2 e Best Effort) e viabilizado para 3 cenários, o atendimento ao cliente que demande um dos serviços *e-science*, considerando o cliente como um domínio particular interagindo com o domínio da rede do provedor, em um segundo cenário o cliente possuindo seu próprio domínio DCN e no terceiro cenário o cliente participando do domínio DCN do provedor.

Utiliza diferentes tecnologias no plano de dados, além das funcionalidades de monitoramento dos circuitos e da rede. O serviço é de natureza estatística, porém, com garantia de um serviço orientado a conexão. Para o seu estabelecimento é necessário o relacionamento entre diferentes tecnologias e os caminhos fim a fim.

O aplicativo usa um modelo simples de interação com o usuário através de interface *web*, mantendo as características e atributos para cada modalidade do serviço a ser utilizado. O domínio DCN do provedor é o responsável pelo estabelecimento lógico de cada interface requisitada.

Algumas etapas são necessárias na implementação e prestação deste serviço, como: avaliação dos impactos na rede, a escolha do cenário mais adequado a rede do cliente, regras de diagnóstico e recuperação em caso de falhas, monitoramento pró-ativo, adaptação dos sistemas e processos na rede do provedor e cliente.

5.CONCLUSÃO

Esta dissertação apresenta um serviço - Profile Service - para aplicações *e-science* a ser implantado em redes híbridas. É um serviço de comunicação de dados em uma arquitetura de rede com comutação dinâmica. Esta comutação inicia com a solicitação do cliente, interagindo com o plano de controle automático da rede, com a grande vantagem de possibilitar o atendimento das necessidades de cada cliente e durante período especificado. Isto traz para as redes otimização dos recursos, flexibilidade e agilidade no atendimento à demanda, uma vez que seu estabelecimento é dinâmico.

Considerando que:

- Novos cenários de rede são desenvolvidos;
- Toda e qualquer evolução tecnológica deve ser feita de forma cautelosa;
- Os recursos para investimento em novas tecnologias são escassos;

É necessário:

- Aumento da capacidade de oferecimento de novos produtos;
- Aumento da produtividade;
- Redução dos custos;
- Criação de novos serviços.

Nestas novas arquiteturas, os desafios de integrar as diversas plataformas existentes, sem perder os benefícios de cada uma é algo relevante. Uma alternativa, recentemente considerada, é a definição de modelos nos quais arquiteturas dos protocolos e serviços podem ser dinamicamente programados ou adaptados ao longo de seu ciclo de vida. Essa adaptabilidade é uma característica desejável tanto para acomodar e absorver os avanços tecnológicos que acontecem cada vez mais rapidamente, como para permitir a criação de novos serviços e aplicações que apresentem requisitos específicos de QoS.

A necessidade de compartilhar informações entre pessoas tem se tornado tão importante que provoca um aumento significativo do tráfego de dados, voz, vídeo e *Internet*. Este crescimento vem criando um mercado competitivo e a necessidade de novos serviços, levando à integração das redes tradicionais. Porém, hoje estas redes ainda não são assim. Os

maiores desafios que temos que enfrentar é a convivência com novas tecnologias e atendimento aos serviços tradicionais que ainda demandam recursos nas redes das operadoras de telecomunicações, como a rede legada, além dos sistemas de gerência existentes ainda não serem fim a fim na criação de circuitos.

Tradicionalmente as redes de comunicação de dados podem ser divididas em redes que oferecem canais dedicados e as que utilizam técnicas de comutação de pacotes.

As redes híbridas, em especial aquelas com comutação dinâmica, prometem o melhor dos dois mundos com importantes ganhos operacionais. Recentemente, algumas aplicações científicas demandam a transmissão de vídeos com alta resolução e altas capacidades, muitas vezes com aplicações de computação distribuída em *grid* ou acessos a instrumentos remotos. E como as redes vão se comportar com estas novas demandas? Serão necessário aumento das capacidades dos enlaces, garantias de QoS, com isolamento do tráfego e caminhos dedicados, circuitos pré-estabelecidos e aprovisionamento dinâmico, via sinalização ou gerenciamento

Foi identificada também a necessidade de fazer um cuidadoso estudo de capacidade da rede, bem como das suas vulnerabilidades de segurança. Precisam também ser feitos novos estudos, ou mesmo desenvolvimentos, para obter escalabilidade, confiabilidade e disponibilidade da plataforma e do serviço. Todo este trabalho deve ser validado pela realização de um piloto do serviço com a inclusão de usuários experimentais, inicialmente sem a garantia de qualidade ou nível de serviço, antes de sua liberação para comercialização para verificação de escalabilidade.

As redes DCN abrem grandes e inovadoras perspectivas para o mercado de telecomunicações permitindo atender a demandas reprimidas e desenhar novos serviços. No escopo de uma NREN, o serviço DCN permite suportar e mesmo fomentar novas atividades de *e-science*, computação em *grid* e trabalhos colaborativos.

A contribuição desta dissertação é a de especificar as funcionalidades e as necessidades de adaptação de sistemas e processos para entrega do novo serviço Profile Service nas redes dos clientes. Este serviço é implementado de forma transparente, como um serviço de comunicação de dados que utiliza a arquitetura das redes híbridas e comutação dinâmica para prover circuitos ponto a ponto de altas capacidades. A solicitação para estabelecimento de tais circuitos começa a partir do cliente. A proposta de forma organizada e com um plano de planejamento e implementação, descrição das primitivas do serviço e

políticas, que organizam as lógicas do serviço, como implementação, topologia e QoS é descrito neste trabalho.

O serviço foi planejado com 3 modalidades, que possuem conjunto de atributos que definem as suas características (Gold, Silver L2 e Best Effort) e viabilizado para três cenários de atendimento ao cliente que demande um dos serviços *e-science*. Em um primeiro cenário considera-se o cliente como um domínio particular interagindo com o domínio da rede do provedor; em um segundo cenário, o cliente possui o seu próprio domínio DCN; e no terceiro cenário o cliente participando do domínio DCN do provedor.

Para os três cenários descritos poderão ser reavaliados os testes dos serviços, quando as novas placas dos equipamentos, possuírem a interface UNI 2.0 implementada, então poderemos oferecer os serviços dinâmicos de forma automática. Os serviços serão mapeados com definição de QoS em uma fila específica nos roteadores do provedor, de forma a garantir que o tráfego da rede DCN, fique separado do tráfego da rede hoje em operação. Isto fica como sugestão para outros trabalhos.

6.REFERÊNCIAS BIBLIOGRÁFICAS

- [1] AGUIAR, E. C. V. P. Aguiar, “Relatório Técnico e Resenhas Padrões ITU-T e IETF”, Outubro 1999
- [2] AGUIAR, E. C.V.P. Aguiar, "Tecnologias Banda Larga", UFF, Junho 2007
- [3] AGUIAR, E. C.V.P. Aguiar, E. Guia, M., "QoS em Redes sem Fio", UGF, Julho 2005
- [4] BALAKRISHNAN, R. Balakrishnan, “*Advanced QoS for Multi-Service IP/MPLS Networks*”, 2008
- [5] CARNEIRO, L. Ferreira, “Processos Para o Gerenciamento de Rede Híbridas”, projeto dissertação UNIRIO, Setembro 2011
- [6] CHAN, V. CHAN, “*Near-Term Future of the Optical Network in Question?*”, IEEE – *Journal on Selected Areas in Communication*, vol 25, nº9, Dezembro 2007
- [7] CHAN,V.CHAN,“*OpticalFlowSwitching*”,http://www.mit.edu/~medard/papersnew/WOBS_Final, acesso em Dezembro 2010
- [8] CINTRA, C. Cintra, “MPLS-TE”, Junho 2009, projeto final
- [9] CLARK, David D. Clark, “*The end-to-end and Application Design: The role of Trust*”, 1984
- [10] FERREIRA, A. E. Ferreira, Aguiar, E. C. V. P., “Caracterização do serviço de DCN 04_11_10”, outubro de 2010
- [11] FURTADO, C. S. Furtado, "GMPLS", ESTACIO, Julho 2007, projeto final
- [12] G.8080, ITU-T Rec. G.8080/Y.1304, "*Architecture for the Automatically Switched Optical Network (ASON)*", June 2006
- [13] IEEE *Communications Magazine*, D. Simeonidou¹, E. Escalona¹, G. Zervas¹, R. Nejabati¹, S. Spadaro², A. Binczewski³, G. Carrozzo⁴, N. Ciulli⁴, "*A Grid-enabled Control Plane Architecture: The PHOSPHORUS approach*", 2008
- [14] IEEE *Communications Magazine*, H. Sunan "*Designing Hybrid SONET and DWDM Networks*", June 2006
- [15] IEEE *Communications Magazine*, Lehman T., Sobieski J., Jabbari B., "*DRAGON: A Framework for Service Provisioning in Heterogeneous Grid Networks*", June 2006
- [16] IEEE *Communications Magazine*, Okamoto S., Otsuki H., Otani T., "*Multi-ASON and GMPLS Network Domain Interworking Challenges*", June 2008
- [17] IEEE *Communications Magazine*, Sun W., Xie G., Jin Y., Guo W., Hu W., Lin X., and Wu M., "*A Cross-Layer Optical Circuit Provisioning Framework for Data Intensive IP End Hosts*", February 2008
- [18] IEEE *Optical Fiber Communication and the National Fiber Optic Engineers Conference*, Shukia V., Brown D., Hunt C., Mueller T., Varma E., "*Next Generation Optical Network - Enabling Dynamic Bandwidth Services*.4348534 abstract", March 2007
- [19] KUROSE, J. F. Kurose, K. W. Ross, "Redes de Computadores e a Internet", 2006

-
- [20] MINEI, I. Minei, J. Lucek, "*MPLS – Enabled Applications*", October 2005
- [21] MONTEIRO, D. B. Monteiro, "Proposição de Suporte a Serviços Atendidos por DCN", UFF, Dezembro 2011, dissertação
- [22] OIF, "*User Network Interface (UNI) 2.0 Signaling Specification*", <http://oiforum.com>, acessado em Dezembro de 2010
- [23] OSBORNE, E. Osborne, A. Simha, "*Traffic Engineering with MPLS*", Cisco Press, July 2002
- [24] RFC2205, Braden, R., Zhang, L., Berson, S., Herzog, S., and S. Jamin, "*Resource ReSerVation Protocol (RSVP) -- Version 1 Functional Specification*", September 1997
- [25] RFC2328, Moy, J., "*OSPF Version 2*", April 1998
- [26] RFC2702, Awduche, D., Malcolm, J., Agogbua, J., O'Dell, M. and J., McManus, "*Requirements for Traffic Engineering Over MPLS*", September 1999
- [27] RFC3031, Rosen, E., Viswanathan, A., and R. Callon, "*Multiprotocol Label Switching Architecture*", January 2001
- [28] RFC3036, Andersson, L., Doolan, P., Feldman, N., Fredette, A., and B. Thomas, "*LDP Specification*", January 2001
- [29] RFC3213, Ash, J., Girish, M., Gray, E., Jamoussi, B. and G. Wright, "*Applicability Statement for CR-LDP*", January 2002
- [30] RFC3474, Z. Lin, D. Pendarakis, "*Documentation of IANA assignments for Generalized MultiProtocol Label Switching (GMPLS) Resource Reservation Protocol - Traffic Engineering (RSVP-TE) Usage and Extensions for Automatically Switched Optical Network (ASON)*", March 2003
- [31] RFC3630, Katz, D., Kompella, K., and D. Yeung, "*Traffic Engineering (TE) Extensions to OSPF Version 2*", September 2003
- [32] RFC3945, Mannie, E., Ed., "*Generalized Multiprotocol Label Switching (GMPLS) Architecture*", October 2004
- [33] RFC4201, Kompella, K., Rekhter, Y., and L. Berger, "*Link Bundling in MPLS Traffic Engineering (TE)*", October 2005
- [34] RFC4203, Kompella, K. Ed. and Y. Rekhter, Ed., "*OSPF Extensions in Support of Generalized Multi-Protocol Label Switching (GMPLS)*", October 2005
- [35] RFC4208, G. Swallow, J. Drake, H. Ishimatsu, Y. Rekhter, "*Generalized Multiprotocol Label Switching (GMPLS) User-Network Interface (UNI): Resource ReserVation Protocol-Traffic Engineering (RSVP-TE) Support for the Overlay Model*", October 2005
- [36] RFC4209, A. Fredette, Ed., J. Lang, Ed., "*Link Management Protocol (LMP) for Dense Wavelength Division Multiplexing (DWDM) Optical Line Systems*", October 2005
- [37] RFC4258, Brungard, D., "*Requirements for Generalized Multi-Protocol Label Switching (GMPLS) Routing for the Automatically Switched Optical Network (ASON)*", November 2005

-
- [38] RFC4364, E. Rosen, Y. Rekhter " *BGP/MPLS IP Virtual Private Networks (VPNs)*", February 2006
 - [39] RFC4394, Fedyk, D., Aboul-Magd, O., Brungard, D., Lang, J., and D. Papadimitriou, "A *Transport Network View of the Link Management Protocol (LMP)*", February 2006
 - [40] RFC4420, A. Farrel, Ed., D. Papadimitriou, J.-P. Vasseur, A. Ayyangar, "Encoding of Attributes for Multiprotocol Label Switching (MPLS) Label Switched Path (LSP) Establishment Using Resource Reservation Protocol-Traffic Engineering (RSVP-TE)", February 2006
 - [41] RFC4652, D. Papadimitriou, Ed., L. Ong, J. Sadler, S. Shew, D. Ward, "Evaluation of Existing Routing Protocols against Automatic Switched Optical Network (ASON) Routing Requirements", October 2006
 - [42] RFC5654, B. Niven-Jenkins, D. Brungard, M. Betts, "Requirements of an MPLS Transport Profile", September 2009
 - [43] SITE CPqD, <http://www.cpqd.com.br/pad-e-inovacao/272-plataformas-ip/2376-redes-convergentes.html>, Fevereiro 2013
 - [44] SITE DRAGON, <http://dragon.maxgigapop.net>, "Policy-Based Resource Management and Service Provisioning in GMPL", Setembro 2010
 - [45] SITE GEANT, <http://www.geant2.net/>, acesso em Dezembro de 2010
 - [46] SITE IEEE, <http://www.ieee802.org/1/>, acesso em Julho de 2011
 - [47] SITE IETF, <http://tools.ietf.org/html/>, acesso em Julho de 2011
 - [48] SITE INTERNET2, <http://www.internet2.edu/>, "Internet2 Interoperable On-demand Network (ION)", acesso em Dezembro de 2010
 - [49] SITE IPOP2010, <http://www.pilab.jp/ipop2010/info/onlineproceedings.html>, acesso em Dezembro 2010
 - [50] SITE ITU-T, <http://www.itu.int/itu-t/recommendations/index.aspx?ser=G>, acesso em Julho de 2011
 - [51] SITE MEF, <http://metroethernetforum.org/index.php>, "Ethernet Services Model", acesso em Setembro de 2011
 - [52] SITE OIFORUM, www.oiforum.com/public/downloads/oif-p0040%20042%2000-HW.pdf, "On-Demand Ethernet Services", acesso em Janeiro de 2011
 - [53] SITE ONTC, <http://www.ontc-prism.org/>, Ambrosia, J. D, "The Next Generation of Ethernet", acesso em Dezembro de 2010
 - [54] SITE OSCARS, <http://www.oscars.es.net/OSCARS/docs/>, "Guok C., "ESnet On-demand Secure Circuits and Advance Reservation System", acesso em Julho de 2011
 - [55] SPADARO, D. Simeonidou¹, E. Escalona¹, G. Zervas¹, R. Nejabati¹, S. Spadaro², A. Binczewski³, G. Carrozzo⁴, N. Ciulli⁴, "A Grid-enabled Control Plane Architecture", http://upcommons.upc.edu/e-prints/bitstream/2117/9383/1/OFC_Spadaro.pdf, acessado em Dezembro 2010
 - [56] The IDCP Control Plane Web Site, <http://www.controlplane.net>, September 2011
 - [57] The PHOSPHORUS, <http://dragon.maxgigapop.net>, September 2010

-
- [58] T-REC-G.694.1, ITU-T “*Spectral grids for WDM applications:DWDM frequency grid*”, 2009
 - [59] T-REC-G.707, ITU-T “*Network node interface for the synchronous digital hierarchy (SDH)*”, 2007
 - [60] T-REC-G.709, ITU-T “*Interfaces for the Optical Transport Network (OTN)*”, 2009
 - [61] T-REC-G.841, ITU-T “*Types and characteristics of SDH network protection architectures*”, Outubro 1998
 - [62] T-REC-G.872, ITU-T “*Architecture of optical transport networks*”, 2010
 - [63] T-REC-G.992, ITU-T “*Asymmetric Digital Subscriber Line (ADSL)*”, 2003
 - [64] T-REC-I.200, ITU-T “*Recommendation I.200, Integrated Service Digital network Service Capabilities*”, 1993
 - [65] WIKI - http://docwiki.cisco.com/wiki/MPLS/Tag_Switching, acesso fevereiro 2013
 - [66] WIKI - http://en.wikipedia.org/wiki/Ipsilon_Networks, acesso fevereiro 2013
 - [67] WIKI - <http://pt.wikipidia.org/wiki/Drupal>, acesso novembro 2012
 - [68] WIKI - Wiki do projeto RNP - <http://wiki.rnp.br/display/futura/RedeCipo>, acesso fevereiro 2010
 - [69] WIKI- NSF- http://www.nsf.gov/funding/pgm_summ.jsp?pims_id=12767, acesso fevereiro 2013
 - [70] WIKI-DCN/IONhttp://groups.geni.net/geni/attachment/wiki/GpENI/ION_in_GpE..., acesso fevereiro 2013
 - [71] YANG, Xi, Lehman T., Tracy C., Sobieski J., Gong S., Torab P., Jabbari B., “*Policy_Based Resource Management and Service Provisioning in GMPLS Network*”, IINFOCOM April 2006
 - [72] YOUNG, G. Young, “*Optical Transport Networks & Technologies Standardization Work Plan Issue 12*”, October 2009