

Este exemplar corresponde a redação final da tese
defendida por Carlos Eduardo
Pagani e aprovada pela Comissão
Julgada em 22/05/20
Maurício F. Magalhães
Orientador

**Proposta e Análise de uma Arquitetura de
Serviços Integrados IP sobre ATM**

UNICAMP
BIBLIOTECA CENTRAL
SEÇÃO CIRCULANTE

Autor: Eng. Carlos Eduardo Pagani
Orientador: Prof. Dr. Maurício F. Magalhães
Examinador: Prof. Dr. Sibelius Lellis Vieira
Examinador: Prof. Dr. Eleri Cardozo
Examinador: Leonardo de Souza Mendes

Dissertação submetida à Faculdade de Engenharia
Elétrica e de Computação da Universidade Estadual de
Campinas como parte dos requisitos exigidos para a
obtenção do Título de Mestre em Engenharia Elétrica.

Dissertação de Mestrado

Campinas / SP - Brasil
Maio 2000

UNIDADE	BC
N.º CHAMADA:	TIUNICAMP
	P14p
V.	Ex.
TOMBO BC/	45296
PROG.	16-392101
C	☐
D	☒
PREC.	R\$ 11,00
DATA	24/07/01
N.º CPD	

CM00157770-9

FICHA CATALOGRÁFICA ELABORADA PELA
BIBLIOTECA DA ÁREA DE ENGENHARIA - BAE - UNICAMP

P14p

Pagani, Carlos Eduardo

Proposta e análise de uma arquitetura de serviços integrados IP sobre ATM / Carlos Eduardo Pagani.-- Campinas, SP: [s.n.], 2000.

Orientador: Maurício F. Magalhães.

Dissertação (mestrado) - Universidade Estadual de Campinas, Faculdade de Engenharia Elétrica e de Computação.

1. Sistemas de transmissão de dados. 2. Sistemas de telecomunicação. 3. TCP-IP (Protocolo de rede de computação. I. Magalhães, Maurício F. II. Universidade Estadual de Campinas. Faculdade de Engenharia Elétrica e de Computação. III. Título.

RESUMO

Novas aplicações estão surgindo para integrar o uso de dados, voz e vídeo na Internet. Essas aplicações requisitam de qualidade de serviço (QoS - *Quality of Service*) na transmissão dessas informações. Isso motivou o desenvolvimento de serviços adicionais à arquitetura TCP/IP (*Transfer Control Protocol / Internetworking Protocol*) para garantir QoS. Em paralelo, a tecnologia ATM (*Asynchronous Transfer Mode*) é um importante meio de transmissão em redes de alta velocidade devido a capacidade de transmitir dados, voz e vídeo integrados com garantia de QoS. Esta tese aborda as questões referentes ao uso da tecnologia ATM para suportar os novos serviços IP. Neste trabalho um modelo de comutação e integração IP/ATM é proposto e avaliado por meio de simulação de algumas de suas características básicas.

ABSTRACT

New Internet applications integrates data, voice and video forwarding. This applications need Quality of Service (QoS) to control their traffic. Additional services in the TCP/IP (*Transfer Control Protocol / Internetworking Protocol*) architecture were proposed to do this. Additionally, ATM is an important transmission mode in high speed networks due to its capability to send integrated data, voice, and video with QoS. This thesis reports the use of ATM technology to support new IP services. This work proposes a new IP/ATM switching model and simulates the basic characteristics of this architecture.

Luciana,

Sem você estas palavras não estariam sendo lidas...

AGRADECIMENTOS

Ao ensino público brasileiro em todos os níveis, gratuito e de boa qualidade, pois sem ele este trabalho não teria sido realizado.

À CAPES pelo suporte financeiro.

Ao Prof. Dr. Maurício Ferreira Magalhães pelas dicas e conselhos que me levaram a trilhar todo o trabalho de desenvolvimento contido neste trabalho.

Aos colegas da pós-graduação: Alencar, pela indicação de artigos e do caminho inicial; Joaquim, pelo suporte remoto e discussões filosóficas; e a todos que contribuíram de algum modo e, por razões de espaço e tempo, não puderam ser citados.

ÍNDICE ANALÍTICO

RESUMO.....	iii
ABSTRACT.....	iii
ÍNDICE ANALÍTICO.....	vii
LISTA DE FIGURAS.....	x
LISTA DE TABELAS.....	xi
CAPÍTULO 1 - INTRODUÇÃO GERAL	1
1.1 Contexto e motivação.....	1
1.2 Desenvolvimento do trabalho.....	3
1.3 Organização da monografia.....	5
CAPÍTULO 2 - MODELOS DE SERVIÇOS IP	7
2.1 Introdução.....	7
2.2 Modelos de Serviços	7
2.2.1 Mecanismos Gerais para Compartilhamento de Recursos.....	8
2.2.2 Modelo de Serviços Integrados à Internet.....	11
2.2.3 Modelo de Serviços Diferenciados na Internet.....	14
2.3 Protocolo de Reserva de Recursos	15
2.3.1 Características Gerais	17
2.3.2 Modelo de Reservas.....	18
2.3.2.1 Tipos de mensagens	19
2.3.2.2 Estilos de Reserva	22
2.3.3 Criação de uma sessão.....	23
2.3.3.1 Estilos e Fusões de Reservas.....	26
2.4 Modelos de Classes de Serviço	27
2.4.1 Serviço Garantido	28
2.4.1.1 Mensagens e Policiamento	29
2.4.2 Serviço de Carga Controlada	31
2.4.2.1 Mensagens e Policiamento	32
2.5 Conclusão.....	32
CAPÍTULO 3 - SERVIÇOS INTEGRADOS SOBRE ATM	34
3.1 Introdução.....	34
3.2 Modelos IP sobre ATM	35

5.6 Comutação Adaptativa com Fluxos Agregados	94
5.6.1 Agregação de Fluxos.....	95
5.6.2 Roteamento sem Classes de Endereços	97
5.6.3 Tipo de Fluxo com Agregação por Subredes	98
5.6.4 Agregação de tráfego em canais virtuais ATM	100
5.6.5 Alteração no Tipo de Fluxo.....	101
5.6.6 Análise da Agregação de Fluxos.....	103
5.6.6.1 Resultados Obtidos	104
5.7 Conclusões	108
 CAPÍTULO 6 - CONCLUSÃO GERAL	 109
6.1 Considerações Finais	109
6.2 Cenário de Utilização.....	111
6.3 Contribuições e Trabalhos Futuros	113
 REFERÊNCIAS BIBLIOGRÁFICAS	 115
 APÊNDICE A - INTERCONEXÃO DE REDES ATM	 121
A.1 Introdução	121
A.2 Emulação de Rede Local	122
A.2.1 Componentes e Tipos de Conexões.....	122
A.2.2 Operação.....	124
A.3 IP Clássico sobre ATM.....	125
A.3.1 Encapsulamento de Pacotes	125
A.3.2 Subredes Lógicas.....	126
A.4 Modelo baseado em nuvens ATM	127
A.4.1 Redes com Acesso Múltiplo sem Difusão.....	127
A.4.2 Servidores de Difusão	129
A.5 Multiprotocolo sobre ATM.....	130
 APÊNDICE B - ESTRUTURA DE SIMULAÇÃO.....	 132
B.1 Introdução	132
B.2 Arquitetura de simulação	132
B.3 Arquivos de Tráfego	136
B.4 Modificações Implementadas no Simulador	139

LISTA DE FIGURAS

Figura 2.1: Modelo de Referência para Roteadores.....	12
Figura 2.2: RSVP em Host e Roteadores.....	16
Figura 2.3: Fluxo de Mensagens do RSVP.....	21
Figura 2.4: Procedimentos de inserção em uma sessão RSVP	25
Figura 2.5: Exemplos dos estilos de reserva definidos pelo RSVP	26
Figura 3.1: Classificação dos modelos de adaptação ATM com a camada de rede	35
Figura 3.2: Relacionamento entre o modelo IIS e o padrão ATM.....	37
Figura 3.3: Processo de estabelecimento de conexões RSVP em ATM.....	41
Figura 3.4: Múltiplos níveis de QoS para um mesma sessão RSVP	43
Figura 3.5: Solução heterogênea completa	44
Figura 3.6: Solução heterogênea limitada.....	45
Figura 4.1: Rede de comutadores por rótulos de roteamento	55
Figura 4.2: Tipos de Comutação	58
Figura 4.3: Arquitetura <i>IP Switch</i>	62
Figura 4.4: Isolamento de fluxo para comutação.....	64
Figura 4.5: Formas de Preenchimento da TIB	69
Figura 5.1: Arquitetura IS em um comutador por rótulos	79
Figura 5.2: Porcentagem de comutação × VCs disponíveis ($X=5$ $Y=40$ $FTO=90$).....	85
Figura 5.3: Efeito da variação do tempo de desativação de conexões (FTO).....	87
Figura 5.4: Número de conexões ativadas e desativadas por segundo (FixWest1).....	88
Figura 5.5: Estrutura do objeto para distribuição de rótulos.....	92
Figura 5.6: Escalonamento integrado nas camadas de rede e de enlace.....	93
Figura 5.7: Orientação baseada em tráfego × topologia	96
Figura 5.8: Endereçamento CIDR.....	97
Figura 5.9: Definição do tipo de fluxo FT3	99
Figura 5.10: Intercalação de células e filas para reordenação	100
Figura 5.11: Mudança consistente de política de agregação.....	103
Figura 5.12: Comparação do uso dos tipos de fluxos FT2 e FT3	105

Capítulo 1

Introdução Geral

1.1 Contexto e motivação

Atualmente, redes de comunicações específicas e especializadas em um único tipo de serviço estão dando lugar a sistemas integrados para envio de informações. Esses novos sistemas são capazes de oferecer serviços de áudio, vídeo e dados em uma única rede que dá suporte as diversas características de cada tipo de tráfego. Dentro desta tendência, a comunidade internacional de telecomunicações padronizou o conceito de RDSI/FL – Rede Digital de Serviços Integrados/Faixa Larga (B/ISDN – *Broadband/Integrated Services Digital Networks*) como a rede capaz de transportar diferentes tipos de tráfego em um único meio de transmissão.

No final da década de 1980, a tecnologia ATM (*Asynchronous Transfer Mode*) foi escolhida e adotada pela indústria (representada pelo ATM Fórum) como a forma de implementar a infra-estrutura necessária para uma RDSI/FL. Em 1987, o organismo internacional de padronização ITU-T (*International Telecommunications Union – Telecomm standardization sector*) definiu o ATM como a técnica de transferência (meio de transmissão) a ser utilizado pelas redes RDSI/FL. Esses fatos transformaram a tecnologia ATM no padrão mais importante e amplamente aceito para integrar em um só meio o envio de informações de áudio, vídeo e dados [51].

Com essa expectativa de integração das redes telefônicas e de dados em uma única rede de alta velocidade, começaram a ser desenvolvidas novas aplicações que se utilizam intensivamente de recursos multimídia e em alguns casos até características de sistemas em tempo real. Dentre essas novas aplicações podemos citar: vídeo conferência, vídeo sob demanda, voz sobre pacotes, TV de alta definição, redes de dados móveis, entre outras. Essas aplicações demandam uma crescente capacidade de tráfego, implicando em um consumo excessivo da largura de banda, um recurso escasso que deve ser compartilhado por todas as aplicações.

Nesse contexto, a implantação de novos serviços para a arquitetura IP podem se utilizar das facilidades de QoS oferecidas pelas redes ATM como meio de enlace segundo o modelo OSI (*Open Systems Interconnection*). Desse modo, quando um determinado aplicativo requisitar serviços com algum tipo de qualidade, os nós IP devem estar aptos a trocar informações de controle de modo a estabelecer conexões ATM com recursos adequadamente reservados de forma a atender aquela requisição. Essa interação implica na tradução dos parâmetros de QoS e caracterização de tráfego em sinalização ATM adequada e no gerenciamento de conexões IP em canais virtuais.

Entretanto, a interação entre IP e ATM é potencialmente complexa. Existem diferenças de operação entre as duas arquiteturas, principalmente no modo de oferecer QoS, ou seja, em seus modelos de serviço. Esse problema motivou a elaboração deste trabalho que realizou um estudo teórico inicial sobre os possíveis modelos de QoS para a Internet e as soluções existentes para integração com a tecnologia ATM. A partir desse estudo, um novo mecanismo para integrar fortemente as arquiteturas IP e ATM foi proposta e simulada.

1.2 Desenvolvimento do trabalho

O objetivo deste trabalho foi abordar as principais questões relativas a QoS na Internet e sua integração pelas redes ATM. Com esse propósito, uma primeira análise determinaria qual modelo de serviços seria utilizado para habilitar a arquitetura IP a trabalhar com QoS em redes ATM.

No esforço de proporcionar QoS na Internet, dois modelos de serviços distintos, e de certa forma complementares, foram desenvolvidos no IETF: o modelo de Serviços Diferenciados (*Diffserv – Differentiated Services*) e o modelo de Serviços Integrados (*Intserv – Integrated Services*).

O modelo *DiffServ* define grupos de QoS onde os pacotes são classificados de acordo com algum critério. Em um mesmo grupo, os pacotes recebem tratamento comum, sem diferenciação de serviço. Devido ao tratamento comum, o controle de QoS é apenas

O estabelecimento de conexões em comutadores por rótulos pode ser realizada por três modos: (1) orientado por tráfego ou fluxos de dados; (2) orientado por topologia e (3) orientado por mensagens de controle [24].

Como resultado dos estudos realizados, este trabalho propõe uma nova arquitetura de comutação para integrar fortemente os mecanismos do modelo *IntServ* com a tecnologia ATM. Essa nova arquitetura – Comutação IS (*IntServ*) – utiliza o protocolo GSMP (*General Switch Management Protocol*) padrão IETF para controle do comutador que torna a implementação independente de *hardware*. A comutação é inteligente quanto ao estabelecimento de conexões e explora as características de cada modo. A Comutação IS é orientada por fluxos agregados por topologia e permite o estabelecimento de conexões através de mensagens de controle RSVP.

Análises obtidas por simulação das arquiteturas estudadas foram realizadas e resultaram na elaboração da nova proposta buscando otimizar a comutação integrada dos modelos de serviço das arquiteturas IP e ATM. Para validação dessa nova arquitetura, foram realizadas simulações e comparações com os resultados obtidos na simulação das arquiteturas originais propostas na literatura. Por fim, os resultados da simulação são apresentados e discutidos para demonstrar a viabilidade da arquitetura proposta.

1.3 Organização da monografia

Esta dissertação de mestrado está dividida em seis capítulos cujos conteúdos são organizados da seguinte forma:

- Capítulo 1: Introdução Geral – Introduz os problemas abordados neste trabalho. Explica como as abordagens desse trabalho foram desenvolvidas e os motivos de escolha dentre as alternativas que poderiam ser adotadas.
- Capítulo 2: Modelos de Serviços IP – Apresenta os modelos de serviços *IntServ* e *DiffServ* desenvolvidos pelo IETF e seus mecanismos para oferecer QoS na Internet. Exemplifica o funcionamento do protocolo de sinalização RSVP e de classes de serviço adicionais à arquitetura IP.

Capítulo 2

Modelos de Serviços IP

2.1 Introdução

Um modelo de serviços voltado para oferecer QoS consiste de um conjunto de mecanismos que atuam na rede com o objetivo de cumprir os compromissos de QoS firmados na oferta de algum tipo de serviço para um fluxo de transmissão. Esses compromissos são resultantes de requisições provenientes de aplicativos que utilizam a rede para comunicação. Um modelo de serviços para uma rede de pacotes deve levar em consideração questões relevantes da teoria de sistemas de tempo real na definição de seus mecanismos [1].

O termo “Serviços Integrados à Internet” (IIS – *Internet Integrated Services*) [2] é utilizado para definir um modelo de serviços específico que agrega o suporte a características de tempo real à Internet. Este modelo adiciona novos serviços para tratamento do tráfego de dados na arquitetura IP, além de manter o atual serviço de envio sem confiabilidade. O termo “Serviços Diferenciados” (DS – *Differentiated Service*) é utilizado em [3] para definir um modelo com controle de QoS menos rigoroso baseado em grupos de serviço.

A seção 2.2 discute as questões teóricas relevantes aos mecanismos de um modelo de serviços e apresenta um guia para a implementação de tais mecanismos. O modelo IIS e o modelo DS são apresentados. Na seção 2.3 é apresentado o protocolo proposto como mecanismo de configuração de reservas, cuja função é essencial para o modelo de referência IIS. Por fim, a seção 2.4 discute duas classes de serviço propostas para suporte ao controle de QoS para aplicativos.

2.2 Modelos de Serviços

A essência para o suporte a serviços de tempo real é garantir que os requisitos de tempo associados a cada aplicativo sejam satisfeitos. Como um ponto de vista geral

conjunto de estados de reservas em cada elemento de rede (roteador) ao longo de todo o caminho de transmissão. Com o estabelecimento desses estados, uma rota com recursos reservados está disponível na rede para transmissão com garantias de serviço. Essa rota se assemelha às conexões estabelecidas em redes comutadas por circuitos.

Como a simplicidade e a robustez dos protocolos IP é devida à sua natureza não orientada para conexão, optou-se pelo estabelecimento dos estados de reserva de maneira dinâmica, isto é, os estados devem ser renovados periodicamente. Desta forma, em caso de falhas nas rotas, as informações de reserva podem ser restabelecidas sem a necessidade de gerenciamento explícito como feito em conexões. Esta abordagem dinâmica para o estabelecimento de reservas é denominada *soft-state* [2].

Para que as reservas sejam estabelecidas e renovadas, a caracterização de tráfego e a especificação de QoS solicitada pelo fluxo devem ser informadas aos elementos de rede através de um mecanismo explícito de configuração de reservas. As informações enviadas por esse mecanismo são utilizadas para manter registros de utilização de recursos por fluxos nos estados de reserva.

Em um modelo de reserva simples, os estados de reserva na rede não são conhecidos pelo receptores que podem fazer requisições não suportadas pela rede. Para informar o estado de utilização da rede e propiciar requisições mais ajustadas, o mecanismo de configuração implementa o modelo de reserva em dois passos OPWA (*One Pass With Advertising*) [4]. No primeiro passo, uma mensagem especifica o nível de QoS que uma fonte pode oferecer a um receptor, bem como o estado dos enlaces que serão utilizados na transmissão. No segundo, de posse das informações sobre o estado da rede, o receptor pode requisitar um nível de QoS para seus fluxos de acordo com as condições oferecidas pela rede.

Como diversos fluxos requisitam reservas de forma concorrente, deve ser implementado um mecanismo de controle de admissão que execute a atribuição e a gerência dos recursos disponíveis. Esse mecanismo deve avaliar se a requisição de reserva de um fluxo pode ser atendida de acordo com a capacidade da rede. Um novo fluxo é admitido desde que não degrade os serviços já firmados com outros fluxos.

Para que os compromissos de QoS dos diversos fluxos sejam respeitados, deve ser definido um mecanismo para o escalonamento dos pacotes das diversas classes de serviços. Esse mecanismo executa operações de reordenação e descarte dos pacotes de acordo com seu nível de prioridade. Com essas operações, os fluxos recebem um tratamento de acordo com as especificações de garantia de sua classe de serviço.

Além dessas funções, o modelo deve possuir um mecanismo estatístico para monitorar a interface de saída. Esse mecanismo fornece dados sobre os níveis de utilização e disponibilidade dos recursos do enlace. Esses dados e os estados das reservas são utilizados para coordenar o funcionamento dos mecanismos de controle de admissão, de classificação e de escalonamento, sendo armazenados em bases de dados utilizadas pelo modelo de serviço.

2.2.2 Modelo de Serviços Integrados à Internet

Na versão atual do protocolo IP todos os pacotes são enviados por um mesmo tipo de serviço, referenciado na literatura como serviço de melhor esforço (*best-effort*) [5]. Esse serviço usa um esquema de filas baseado unicamente no modelo FIFO (*first-in first-out*) para enviar todos os pacotes quando possível sem nenhum controle de QoS. No modelo IIS, de acordo com classes de serviço, um elemento de rede classifica localmente os pacotes dividindo-os em filas onde podem ser reordenados e reenviados de maneira que a QoS seja mantida para os fluxos prioritários.

O modelo de referência IIS [2] implementa um modelo de serviços composto por quatro componentes principais (utilizando os mecanismos discutidos na teoria da subseção anterior): (1) escalonamento de pacotes; (2) controle de admissão; (3) classificação; e (4) o protocolo de configuração de reservas. A Figura 2.1 mostra como esses componentes se relacionam em um roteador que suporta controle de QoS.

tráfego são usadas para controlar o caminho de dados na interface física segundo um algoritmo de escalonamento.

O conjunto de mecanismos que gerencia as diferentes classes de serviço no caminho de dados é chamado de controle de tráfego, sendo constituído por três componentes do modelo de referência: (1) o controle de admissão, (2) a classificação, e (3) o escalonamento. O funcionamento do controle de tráfego é regido pelos estados armazenados localmente na base de controle de tráfego. Abaixo são especificadas as principais características que esses componentes devem apresentar segundo o modelo de referência.

- **Escalonamento de Pacotes:** gerencia a transmissão dos pacotes usando um conjunto de filas e dispositivos de temporização para reordenar os pacotes nas filas de acordo com suas prioridades. Em situações de sobrecarga, pacotes que possuem baixa prioridade são descartados. O descarte contribui para a redução no tempo de permanência nas filas dos pacotes de alta prioridade.
- **Classificação:** divide os pacotes recebidos em filas de acordo com suas classes de serviço para posterior escalonamento com prioridades. A escolha de uma classe baseia-se na comparação dos dados enviados no cabeçalho do pacote com os parâmetros determinados pela política local. Em uma mesma fila todos os pacotes devem receber o mesmo tratamento do escalonamento de pacotes. Como as classes são definidas de forma local em cada roteador, um mesmo pacote pode ser classificado em diferentes filas por diferentes roteadores ao longo do caminho.
- **Controle de admissão:** o controle de admissão implementa o algoritmo de decisão que o controle de tráfego usa para determinar, de acordo com uma política local, se concede a um fluxo uma determinada QoS requisitada. Um fluxo só é admitido caso existam recursos disponíveis para seu tratamento, sem que sejam degradados os serviços já oferecidos a outros fluxos. O controle de admissão é algumas vezes confundido com policiamento, porém, esse último tem como função principal garantir que um fluxo não desrespeite as características de tráfego acordadas com a rede [2].

A função de classificação atribui aos pacotes de um fluxo um código PHB (*Per Hop Behavior*) [3,41,42] que descreve qual tratamento o fluxo vai receber. Cada grupo PHB é identificado por um código distinto encapsulado no campo DS (*Differentiated Service*). O campo DS foi criado pelo modelo *DiffServ* e utiliza o mesmo espaço destinado ao campo TOS. O campo DS mantém compatibilidade com o uso inicial do campo TOS e todos os códigos já definidos na versão IP 4.

O PHB descreve de forma genérica o comportamento dos pacotes na rede. Fica a cargo de cada implementação a especificação dos algoritmos de escalonamento e gerenciamento de filas para formar o PHB. Os roteadores alocam recursos (como largura de banda e espaço nas filas) para os fluxos de acordo com o PHB específico.

A classificação pode ser realizada em dois modos: (1) modo BA (*Behavior Aggregate*), onde apenas o campo DS é utilizado para classificação; este modo é também conhecido como modo agregado, pois diferentes fluxos podem ter a mesma classificação; ou (2) modo MF (*Multi-Field*), onde mais campos do cabeçalho são utilizados. A classificação deve estar de acordo com o nível de serviço (SLA – *Service Level Agreement*) acordado com o fluxo e deve executar a função de condicionamento de tráfego (TCA - *Traffic Conditioning Agreement*) sobre os pacotes de um fluxo.

2.3 Protocolo de Reserva de Recursos

O protocolo RSVP (*Resource ReSerVation Protocol*), proposto e padronizado pelo IETF¹ [6-8], é um mecanismo de configuração de reservas projetado segundo as características do modelo IIS para atuar como agente de reserva.

Um aplicativo utiliza o protocolo RSVP para requisitar à rede um serviço de transmissão com QoS apropriada. O RSVP é utilizado por um roteador para receber as requisições de serviços dos diversos aplicativos através da rede. As informações

¹O RSVP é resultado do esforço de trabalho de dois grupos do IETF: *IntServ WG (Working Group)* e o *RSVP WG*.

O protocolo RSVP transporta informações sobre encaminhamento de pacotes, caracterização de tráfego e especificação de QoS dos fluxos, entre outras, através de mensagens encapsuladas em objetos. O protocolo define mecanismos especiais para suportar de maneira eficiente a transmissão de dados entre diferentes fontes para diversos receptores. As próximas subseções abordam as características gerais e apresentam os mecanismos específicos utilizados pelo protocolo.

2.3.1 Características Gerais

O protocolo RSVP trabalha no modo de operação simplex, isto é, as reservas de recursos são estabelecidas em apenas uma direção, possibilitando transmissão de dados em apenas um sentido (tanto no caso *unicast* como no *multicast*). O protocolo possui a característica de ser orientado ao receptor, ou seja, o receptor é encarregado de requisitar as reservas no sentido inverso do tráfego dos pacotes de dados. Isto porque, de acordo com seus projetistas [6], a fonte seria sobrecarregada se cuidasse das reservas de todos os receptores e, além disso, o receptor pode especificar melhor as reservas já que tem mais informações sobre as necessidades de seus aplicativos.

O RSVP não é um protocolo de roteamento, apenas interage com os protocolos de roteamento existentes (e até com os futuros) para determinação das rotas necessárias para a transmissão dos pacotes de dados e controle.

Optou-se pela característica *soft-state* no desenvolvimento do RSVP para que o protocolo apresente a mesma qualidade e características tradicionais dos protocolos IP. Desta forma, as informações de controle de um fluxo (tal como os estados de reservas) devem ser atualizadas periodicamente (por *refresh*) através de mensagens enviadas pelas fontes e pelos receptores aos roteadores. Caso contrário, após um determinado período de tempo (*timeout*), as informações são descartadas e as reservas desativadas. Como as mensagens são enviadas utilizando o serviço de transporte IP tradicional, isto é, sem confiabilidade na entrega, o *timeout* das informações deve ser configurado com um valor de tempo necessário para se enviar κ mensagens de *refresh*. Deste modo, mesmo que $\kappa-1$ pacotes sejam perdidos, não haverá desativação das informações de reservas.

2.3.2.1 Tipos de mensagens

O transporte de mensagens RSVP é realizado de maneira opaca, isto é, as mensagens são apenas entregues ao receptor que tem a função de interpretar as informações recebidas. Para tanto, as mensagens são encapsuladas dentro de objetos cujas formas não são definidas e nem conhecidas pelo protocolo. Esses objetos são enviados logo após o cabeçalho RSVP. Deste modo, novas mensagens podem ser definidas na medida da necessidade, sendo suportadas pelo RSVP sem qualquer mudança na especificação do protocolo, já que as informações entregues são avaliadas apenas pelo aplicativo.

As principais mensagens transportadas pelo RSVP são as mensagens de caminho (PATH) e as mensagens de reserva (RESV). As mensagens PATH têm a função principal de estabelecer o caminho para transmissão dos dados desde a fonte até os receptores. O caminho é construído através da utilização de um campo denominado PHOP (*previous-hop*) [8]. Este campo contém o endereço da máquina adjacente que enviou a mensagem PATH, sendo armazenado localmente para que as mensagens de reserva possam ser transmitidas em sentido inverso ao fluxo dos pacotes de dados. As mensagens PATH transmitem também informações sobre o tráfego gerado pela fonte e os níveis de QoS que a rede pode oferecer aos receptores, para que esses montem requisições mais ajustadas à capacidade suportada pela rede, aumentando as chances de um sucesso na admissão (reserva OPWA). As mensagens PATH possuem campos próprios com informações específicas [8] que serão abordados a seguir:

SENDER_TEMPLATE – descreve os pacotes transmitidos pela fonte através dos campos de identificação no cabeçalho (como endereço fonte, identificador de protocolo, portas de entrada e saída). Essa informação é utilizada junto com o objeto *FILTER_SPEC* (definido à frente) para a seleção de alguns pacotes de determinadas fontes de uma sessão.

- Classe de Serviço – define o tipo de serviço requisitado (apresentados na seção 2.4 como garantido ou carga controlada). De acordo com a classe são definidos os mecanismos específicos de controle de fluxo e o tipo de garantia solicitada pelo fluxo, além de especificar os tipos de parâmetros utilizados na caracterização das reservas.
- Especificação de QoS (*RSPEC*) – define a QoS desejada através de parâmetros específicos de cada tipo de classe. Esses parâmetros podem especificar, por exemplo, níveis de atraso ou largura de banda para uma classe de serviço.
- Caracterização de tráfego (*TSPEC*) – descreve o comportamento de tráfego desejado pelo receptor em parâmetros *token-bucket*, ou seja, este objeto especifica a taxa de transmissão e o valor máximo de rajada do tráfego a ser recebido.

FILTER_SPEC – informa aos elementos da rede os pacotes escolhidos para recepção. Esses pacotes são especificados através de alguns campos do identificador de sessão, e que são comparados ao objeto *SENDER_TEMPLATE*.

A Figura 2.3 mostra um exemplo de um fluxo de mensagens em uma sessão RSVP composta de duas fontes (S1 e S2) e dois receptores (R1 e R2) atravessando uma rede formada por três elementos de redes (roteadores).

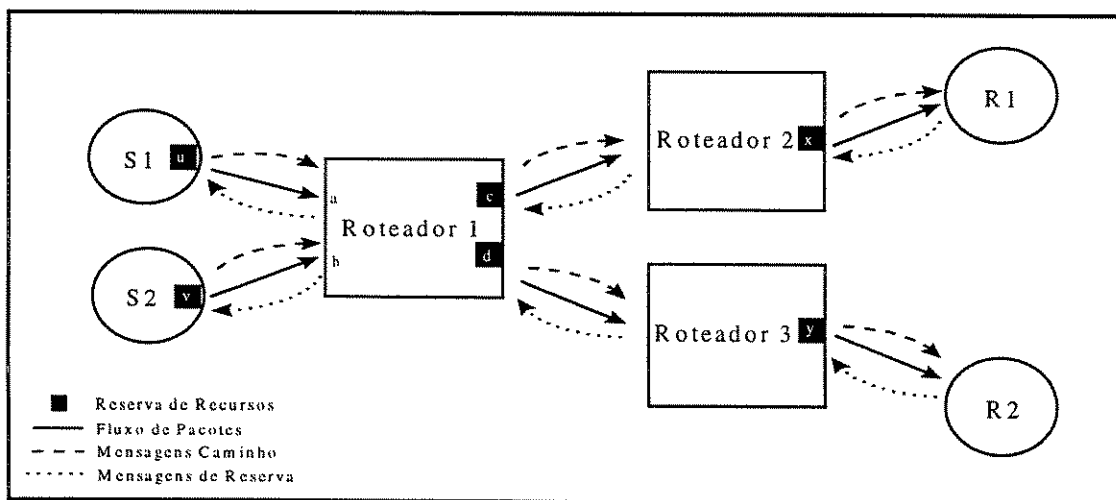


Figura 2.3: Fluxo de Mensagens do RSVP

Estilo WF (*Wildcard-Filter*) – cria um único canal compartilhado por todas as fontes de uma sessão. Notação: WF(*{ Q }), fontes implícitas (“*” significa qualquer) e um nível de QoS Q (definido como um valor abstrato) compartilhado entre todas as fontes no único canal.

Estilo FF (*Filter-Fixed*) – cria um canal exclusivo para cada fonte explicitamente escolhida. Notação: FF($S_1\{Q_1\}$, $S_2\{Q_2\}$, ...), fontes explicitamente escolhidas (S_1 , S_2 , ...) pelo campo *FILTER_SPEC* com seus respectivos níveis de QoS (Q_1 , Q_2 , ...) para cada canal exclusivo.

Estilo SE (*Shared-Explicit*) – cria um único canal compartilhado pelas fontes explicitamente escolhidas para a sessão. Notação: SE((S_1 , S_2 , ...) { Q }), fontes explicitamente escolhidas (S_1 , S_2 , ...) e um único nível de QoS Q compartilhado.

2.3.3 Criação de uma sessão

Para receber dados de uma determinada sessão, um receptor precisa conhecer o identificador de sessão. O meio (automático ou por configuração) pelo qual o receptor toma conhecimento do identificador não é definido pelo modelo IIS. Com o conhecimento do identificador, como mostra a Figura 2.4, seguem-se os seguintes passos para recepção de dados com QoS requisitada:

1. O receptor deve se tornar membro do grupo da sessão através da requisição de entrada a um protocolo de difusão (*multicast*) de mensagens, como por exemplo, um *join-member* dos protocolos ICMP/IGMP (*Internet Control Message Protocol / Internet Group Management Protocol*), respectivamente, para as versões 4 e 6 do protocolo IP.
2. O roteador que recebe a requisição de inserção do receptor inunda a rede com mensagens (pelos mecanismos do protocolo de difusão) para que as tabelas de rotas em todas as máquinas da rede sejam atualizadas.
3. As mensagens PATH começam a chegar ao receptor membro da sessão. Estas mensagens estabelecem o caminho onde as mensagens RESV serão enviadas para

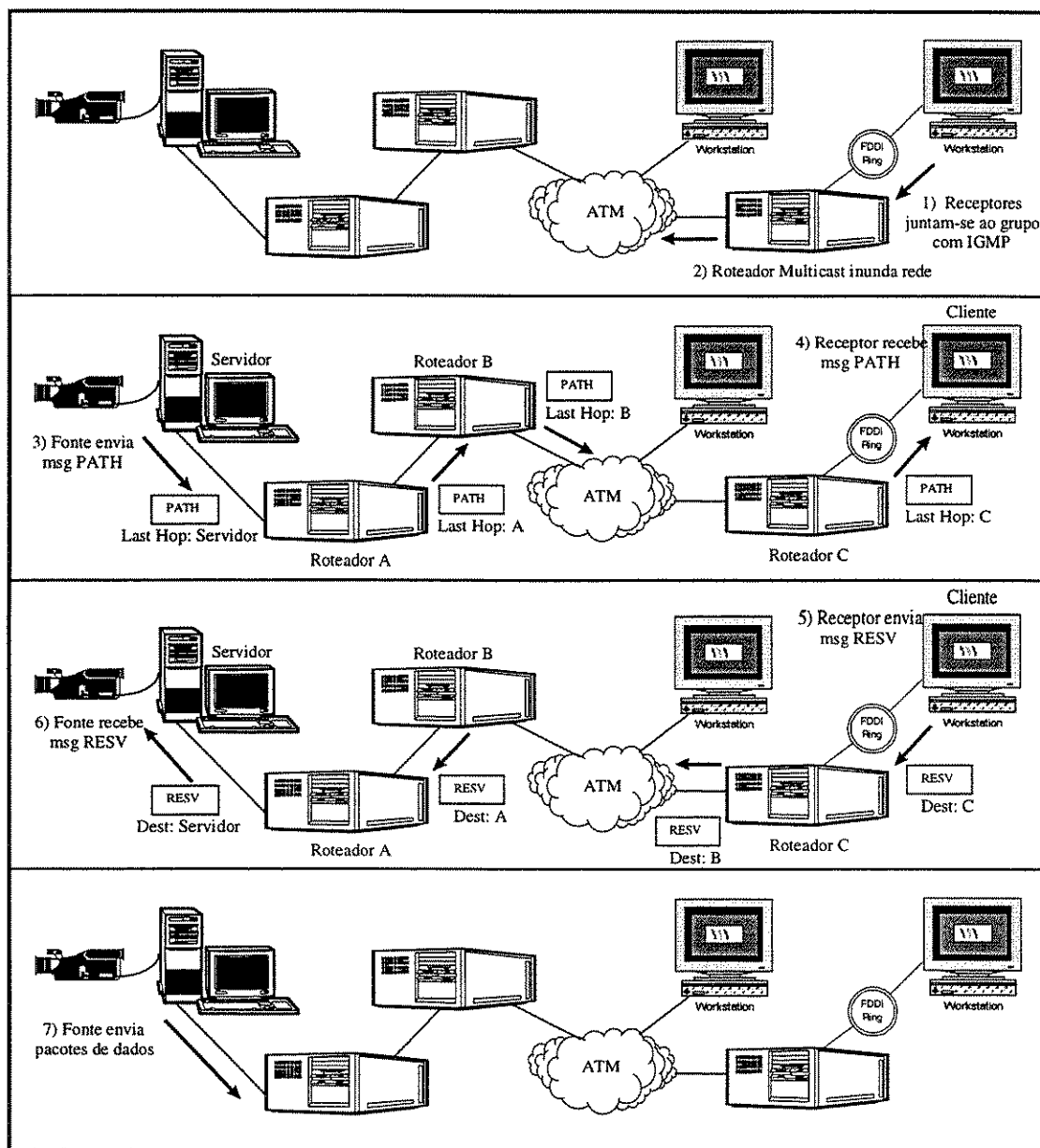


Figura 2.4: Procedimentos de inserção em uma sessão RSVP

A especificação do protocolo RSVP considera que todos os membros de uma sessão devem receber dados, mesmo que não façam reservas. Deste modo, após o passo 1, o receptor começa também a receber pacotes da sessão, através do serviço de melhor esforço. Da mesma forma, mesmo que um receptor pare de enviar mensagens RESV, os dados continuam sendo enviados sem controle de QoS (após o período de *timeout*).

As requisições de reservas atravessam a rede passando por um mesmo elemento de rede (acima na esquerda da figura) onde são descritos os exemplos. O roteador recebe as requisições de três receptores pelas interfaces de saída (c para R1; d para R2 e R3) e as encaminha às fontes através de suas interfaces de entrada (a para S1; b para S2 e S3) executando as ações de reservas e fusões. Assume-se que cada receptor pode receber tráfego proveniente de qualquer uma das fontes.

No estilo WF (abaixo à esquerda), uma única reserva de recurso é estabelecida em cada interface sendo compartilhado por todas as fontes, o maior entre todos os valores de QoS é escolhido como reserva. A mensagem enviada para as fontes é fundida utilizando-se a QoS de maior valor requisitada pelos receptores.

No estilo FF (acima à direita), uma reserva é estabelecida para cada fonte explicitamente citada. Não há fusão na reserva de recursos em cada interface porque os canais são exclusivos. A fusão de mensagens de reservas é ilustrada com o envio do valor $4B$, maior das requisições para a fonte S1. Nas outras fontes não há fusão.

No estilo SE (abaixo à direita), um único canal é estabelecido em cada interface com o maior valor de QoS requisitado entre os receptores. Como o canal é compartilhado, uma única fusão com o maior valor de QoS requisitado é enviada para as fontes pelas interfaces de entrada.

2.4 Modelos de Classes de Serviço

O modelo IIS propõe a utilização de classes de serviços para agregar controle de QoS à arquitetura IP. No entanto, a especificação de classes de serviço é deixada a cargo de propostas independentes sendo apenas definidos aspectos gerais pelo grupo de trabalho *IntServ* do IETF [10]. Deste modo, o modelo IIS é independente de classes particulares, suportando novas definições de acordo com a necessidade dos aplicativos.

Duas principais classes de serviço são especificadas pelo IETF: (1) a classe de serviço garantido e (2) a classe de serviço de carga controlada. As interações dessas classes com o protocolo RSVP estão definidas em um documento separado [9]. Essas classes oferecem aos fluxos: parâmetros específicos para caracterização de QoS,

No termo de erro C são adicionados todos os fatores de atraso relativos à taxa de transmissão dos pacotes, como por exemplo, o atraso devido à divisão dos dados em células ATM. Na fórmula de atraso, esse termo de erro é dividido pela taxa R devido a sua dependência a esse valor.

O termo de erro D representa o pior caso no tempo de espera que um pacote pode sofrer devido a características específicas no envio de dados por um tipo de enlace. Esse termo é geralmente devido à tecnologia utilizada para implementar o enlace do elemento de rede (como o tempo de obtenção de um *slot* de transmissão).

Os termos C e D devem ser informados através de objetos aos elementos de rede pelo protocolo de reservas de recursos. Uma função de composição deve ser aplicada aos termos recebidos e aos termos locais para a determinação das parcelas acumuladas parciais C_{sum} e D_{sum} . O receptor pode calcular o atraso fim a fim máximo experimentado por pacote de um fluxo na rede através dos valores da somatória dos termos de erro ($\sum C_j$ e $\sum D_j$), computados na fórmula de atraso.

2.4.1.1 Mensagens e Policiamento

O serviço garantido é invocado pelo envio de objetos nas mensagens RSVP que caracterizam o tráfego (TSpec) e especificam a QoS (RSpec). Como valores TSpec são utilizados os parâmetros *token-bucket* (r , b), como já visto, e mais os valores de taxa máxima de rajadas injetadas na rede (p), tamanho mínimo de transmissão (m) e tamanho máximo de pacote (M). Este último valor deve ser menor que a MTU (*Maximum Transfer Unit*) mínima da rede para evitar fragmentação.

Como valores RSpec são utilizadas a taxa de serviço R (que deve ser maior que a taxa de geração de fichas) e o termo de relaxamento S . O termo S significa a diferença entre o atraso desejado e o atraso obtido usando-se o nível de reserva R e pode ser utilizado por um elemento de rede para reduzir as reservas de um fluxo. Nenhuma especificação de tamanho de filas é incluída, pois essa informação pode ser obtida por cálculos baseados nos parâmetros recebidos pelo protocolo de reservas.

Todo tráfego não conforme aos parâmetros acordados deve ser marcado e enviado na forma de melhor esforço. Essa penalidade é imposta apenas ao atraso máximo do fluxo não conforme sendo indiferente aos outros fluxos.

2.4.2 Serviço de Carga Controlada

A classe de serviço com carga controlada [13] oferece um serviço que se aproxima do obtido por aplicações que se utilizam do serviço de melhor esforço em condições normais, sem congestionamento. O serviço busca dois objetivos principais, desde que a rede funcione sem falhas:

- Uma porcentagem elevada de pacotes deve ser entregue ao destino com sucesso (a porcentagem de descarte de pacotes deve se aproximar da taxa de erro do meio de transmissão);
- O atraso experimentado pela maioria dos pacotes nos elementos de rede não deve exceder em muito o atraso mínimo experimentado em média por todos os pacotes do fluxo.

Essas definições da especificação são bastante subjetivas. Fica a cargo de cada implementação específica do serviço definir esses parâmetros de perda e atraso.

O termo “tempo de rajada” é definido como o tempo requerido para a transmissão uma quantidade máxima de dados de um fluxo de acordo com uma taxa de serviço definida. Baseado nesta definição, um fluxo com carga controlada deve experimentar pouco ou nenhum atraso médio nas filas em períodos de tempo maiores que o tempo de rajada. Também por essa definição, os fluxos devem sofrer nenhuma ou pouca perda por congestionamento em períodos maiores que o tempo de rajada.

Os elementos de rede podem implementar uma solução estatística sobre o comportamento passado do fluxo para decidir se dispõem de recursos suficientes para atender a novas requisições de serviço. Outro método de implementação é a adoção de técnicas de escalonamento apropriadas para garantir que os fluxos atendidos pelo serviço compartilhem o canal de forma apropriada.

quais todo o tráfego deve ser dividido e manipulado pelos mecanismos de envio e gerenciamento, onde o objetivo das classes de serviço é cumprir os compromissos de QoS assumidos com cada fluxo de acordo com seu tipo e características de tráfego.

O protocolo RSVP é o mecanismo utilizado pelo modelo IIS para troca de informações de controle e estabelecimento de reservas. Em geral, quando um aplicativo deseja obter um nível de serviço específico, ele deve usar uma lista de parâmetros para caracterizar o tráfego e especificar QoS a fim de informar a rede as suas necessidades. Os elementos de rede devem empregar estas informações para aferir seus mecanismos de controle para atender as requisições de cada aplicativo.

A implementação do modelo IIS pode se beneficiar de uma tecnologia de enlace que ofereça meios de controle de QoS. No próximo capítulo será explorada a integração entre o modelo IIS e a tecnologia ATM nos modelos existentes na literatura.

Capítulo 3

Serviços Integrados sobre ATM

3.1 Introdução

As redes ATM estão sendo amplamente utilizadas como uma tecnologia de enlace para transmissão de dados. Como principais fatores dessa utilização estão a alta capacidade de transmissão e o suporte ao controle de QoS. Essas características, se oferecidas por um meio de enlace, facilitam a utilização dos mecanismos de QoS que o modelo IIS disponibiliza. Por essas e outras razões descritas neste capítulo, torna-se justificável um esforço para integração do modelo IIS sobre ATM [14].

Na adaptação do modelo IIS, o problema chave surge nos diferentes modos de sinalização adotados pelo protocolo RSVP e pela arquitetura ATM. Esse problema divide-se em duas partes principais: (1) gerenciamento de conexões e (2) a tradução dos parâmetros de tráfego e de QoS [15]. A tradução de QoS se preocupa com o suporte das classes de serviço do modelo IIS por parâmetros ATM. Enquanto que, o gerenciamento de conexões se concentra em determinar de que forma, em quantos e em quais canais virtuais os fluxos RSVP serão suportados.

Este capítulo apresenta os principais problemas que surgem na adaptação do modelo IIS à arquitetura ATM. A seção 3.2 apresenta uma classificação dos modelos IP e ATM [16]. A seção 3.3 discute aspectos gerais da integração IIS sobre ATM e como os mecanismos RSVP se relacionam com o padrão ATM. A seção 3.4 aborda as soluções apresentadas para o problema de gerenciamento de conexões ATM. Por fim, a seção 3.5 discute brevemente a relação das classes de serviço do modelo IIS com os parâmetros de QoS oferecidos pela arquitetura ATM.

Os modelos sobrepostos consideram a rede ATM como uma grande nuvem correspondente ao nível de enlace do modelo OSI. Essa nuvem sobrepõe (esconde) a topologia da rede na visão dos protocolos de roteamento e, desta forma, todos os membros da rede se consideram vizinhos entre si e precisam de uma malha (parcial ou completa) de conexões para interligação comum.

Essa conectividade com pouca hierarquia pode sobrecarregar a computação de rotas pois essa é diretamente proporcional ao número de vizinhos que cada membro possui. Considerando que existam n membros em uma nuvem, e que cada membro possa se conectar diretamente aos outros $n - 1$ membros como se fossem vizinhos, é necessário um total de $(n \cdot (n - 1) / 2)$ conexões para interligar esses n membros. A malha de conexões consome muitos recursos e aumenta o tempo gasto com a inicialização de rotas e posterior gerência.

Os modelos sobrepostos, que são abordados neste capítulo (veja detalhes destes modelos no apêndice A), apresentam limitações devido ao conceito de nuvem [29]. Essas limitações podem ser evitadas com o uso do modelo de comutação por rótulos que será apresentado no próximo capítulo e que foi adotado como enfoque por este trabalho.

3.3 Relacionamento entre os Modelos de Serviço

A tecnologia ATM foi inicialmente desenvolvida para suportar aplicações da área de telecomunicações (transmissão de voz e vídeo). Por isso, desde o seu começo, o suporte a requisitos de tempo real foi considerado como parte integrante de sua arquitetura. Com o seu desenvolvimento e com a padronização da UNI 4.0 [16] (*User Network Interface*) mais categorias de serviços se tornaram disponíveis, inclusive para o transporte de dados, transformando o ATM em um meio de enlace importante.

Aplicativos de tempo real que operam com o serviço de melhor esforço IP têm reduzida a qualidade na recepção de informações sob condições de congestionamento, e potencialmente, têm um uso ineficiente de banda passante. Para resolver esse problema, o modelo IIS foi desenvolvido para adicionar novos serviços à arquitetura IP, de modo que, mesmo sob efeito de congestionamento, os compromissos de QoS possam ser cumpridos.

conexões para suportar os fluxos de uma sessão, funções equivalentes ao contrato de tráfego ATM e a atribuição de canais virtuais.

O modelo IP Clássico sobre ATM resolve apenas parte da adaptação com o padrão ATM [16]. Esse modelo oferece um meio para resolução de endereços e estabelecimento de conexões, mas não trata do problema de gerência de QoS, e desta forma, não usa de modo eficiente todo potencial ATM. Outro problema resolvido, mas não de forma satisfatória do ponto de vista do modelo IIS, é relativo à transmissão *multicast*. A resolução de endereços *multicast* proposta em MARS (*Multicast Address Resolution Server* – vide apêndice A) suporta a abertura de conexões com QoS em um dos modos de operação, mas não de uma forma eficiente.

Essas soluções transportam pacotes IP em três modos principais: (1) como em uma linha telefônica, ou seja, os pacotes são sobrepostos em conexões ATM com taxa de transmissão constante em circuitos virtuais permanentes com potencial desperdício de banda; (2) de maneira clássica, em conexões com taxa de transmissão variável em circuitos virtuais comutados desprezando características de QoS disponíveis no ATM [18]; ou (3) através de conexões diretas MPOA/NHRP utilizando de forma limitada as propriedades de QoS ATM. Apesar de reduzir o problema a soluções bem conhecidas, nenhum desses modos é satisfatório do ponto de vista integrado de gerência de QoS e de utilização eficiente de banda. Desta forma, o problema remanescente reside na integração do protocolo RSVP com a sinalização ATM, concentrado em seus dois principais módulos: gerenciamento de conexões e caracterização de tráfego e QoS.

3.4 Gerenciamento de Conexões

A tabela 3.1 [17] aborda as principais diferenças existentes na operação de estabelecimento de reservas dos modelos IP e ATM. Os receptores requisitam através de mensagens periódicas a atualização das reservas de recursos. Essas mensagens podem mudar as configurações de QoS a qualquer momento (característica *soft-state*).

As soluções existentes tratam o problema pela abordagem de duas questões principais: (1) Heterogeneidade, especificando como os diferentes níveis de QoS de uma sessão devem ser suportados e quantos canais virtuais são necessários; e (2) Comportamento dinâmico de QoS, que se refere ao comportamento em relação às mudanças nos níveis de QoS requisitados pelos membros do grupo receptor.

Essas soluções se diferenciam em relação ao número de canais virtuais necessários, à necessidade do envio de múltiplas cópias de um mesmo pacote e a flexibilidade em relação aos problemas de heterogeneidade e comportamento dinâmico. Nas próximas subseções as soluções são apresentadas e analisadas.

3.4.1 Conexões de Dados

Os principais esquemas de gerenciamento de conexões de dados são diferenciados pela maneira como os fluxos são agregados em canais virtuais. O modo individual caracteriza as soluções em que uma única sessão RSVP utiliza um ou mais canais virtuais para suportar os vários fluxos com diferentes níveis de QoS. O modo agregado caracteriza as soluções que combinam várias sessões RSVP em um único canal virtual ATM.

O modo individual é de fácil implementação, porém implica em um maior gasto de canais virtuais e recursos, além do aumento do tempo de estabelecimento de conexões. O modo agregado apresenta algumas vantagens, entre elas: elimina o tempo de estabelecimento de conexões e ameniza os problemas de heterogeneidade e comportamento dinâmico [19]. Porém, os benefícios da agregação dependem da complexa escolha de recursos em sessões agregadas. A complexidade desta tarefa é devida às mudanças do perfil de tráfego sustentado. Por esse motivo, ainda não existem soluções satisfatórias para a agregação das reservas nas sessões, o que torna essa área um ponto de alta prioridade para pesquisas [20].

inicia o procedimento de estabelecimento do canal virtual (ponto–multiponto) com parâmetros requisitados para o envio de dados.

A especificação dos procedimentos de sinalização para abertura de canais virtuais para o IP Clássico não suporta o uso de QoS [21,22] e associa temporizadores aos canais virtuais para fechar as conexões ociosas. Essas duas características são indesejáveis para a operação RSVP. No primeiro caso, não há meios de garantir QoS. No segundo, os temporizadores podem gerar instabilidade, pois podem desativar canais em que a sinalização de reserva continue sendo enviada, ativando novamente os canais.

3.4.1.2 Heterogeneidade

Em uma sessão *multicast* RSVP é possível que apenas alguns dos receptores enviem mensagens de reservas enquanto outros receptores continuem a receber tráfego melhor esforço sem qualquer tipo de reserva. Além disso, podem existir vários níveis de QoS requisitados em uma mesma sessão para diferentes conjuntos de fontes e receptores. Estas diferentes requisições implicam em diferentes níveis de recursos alocados nos vários ramos da malha de difusão de uma sessão.

Mesmo com a UNI 4.0, a tecnologia ATM só suporta um único nível de QoS em seus canais virtuais *multicast*. Essa característica resulta em um problema denominado de heterogeneidade. O problema de heterogeneidade provoca redundância nos dados que trafegam pela rede, pois para cada diferente nível de QoS requisitado por um grupo de receptores deve ser aberto um canal virtual distinto.

Na Figura 3.4, a região sombreada representa um caminho comum aos canais virtuais de uma mesma sessão. Por essa região devem ser enviadas cópias de pacotes para cada ramo da malha de receptores (R1, R2 e R3) com QoS diferenciada. Desta forma, esse problema aumenta as chances de congestionamento, pois existe redundância de informação que trafega pela rede.

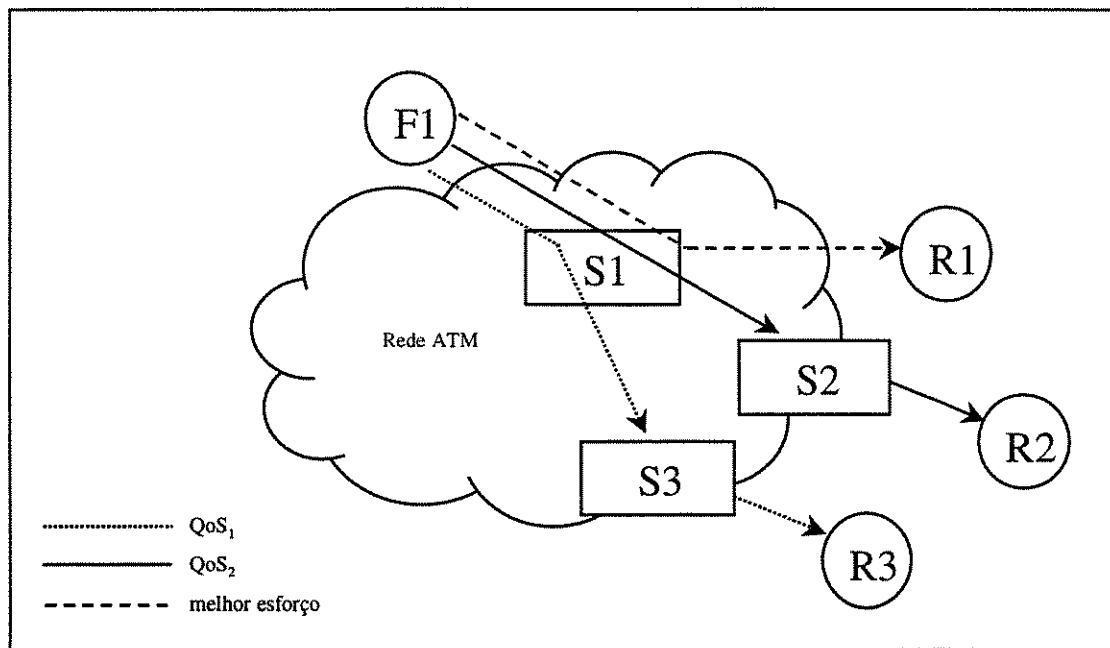


Figura 3.5: Solução heterogênea completa

Solução heterogênea limitada

Limita em dois o número de canais virtuais disponíveis para uma sessão RSVP (Figura 3.6). Um canal é destinado ao tráfego de melhor esforço enquanto o outro é utilizado como única opção de QoS. O nível de recursos alocado para o canal com QoS deve ser suficiente para suportar a maior requisição dos receptores de uma sessão. Desta forma, os recursos alocados também satisfazem todas as outras requisições menores.

Fora da nuvem ATM, as requisições menores são novamente estabelecidas em conexões separadas (como no ramo R3 mostrado na figura). Apesar de diminuir o tráfego em relação à solução anterior, ainda é necessária a duplicação de pacotes nos dois canais. Além disso, essa solução sofre do problema de comportamento dinâmico. Quando surge uma requisição de QoS maior do que a suportada pela conexão, o canal virtual ATM deve ser fechado para que outro que suporte a requisição seja aberto. Esse problema é tratado em detalhes na próxima subseção.

Solução homogênea modificada

Não existe duplicação de dados na solução anterior porém, em alguns casos, alguns receptores podem não receber dados nem na forma de melhor esforço. Esta solução modifica a anterior, permitindo a duplicação de tráfego em caso de rejeição por falta de recursos, estabelecendo um canal de melhor esforço, como especificado na definição do protocolo RSVP. Essa solução, como as duas anteriores, sofre o problema de comportamento dinâmico.

Nenhuma dessas soluções é ótima pois sempre há a possibilidade de existir pelo menos tráfego duplicado na rede. A solução adotada depende do tipo de tráfego suportado. Em uma rede local ATM, a solução heterogênea pode ser suportada, mas em uma rede pública, a solução heterogênea limitada ou a homogênea modificada diminuiria os custos de conexões abertas. Como a especificação RSVP exige que os dados sejam enviados via melhor esforço para os receptores que tenham sofrido falha de admissão de QoS, o modelo homogêneo é o único que não pode ser utilizado.

3.4.1.3 Comportamento dinâmico de QoS

O protocolo RSVP permite que as reservas de recursos possam ser alteradas a qualquer momento da transmissão de uma sessão simplesmente com a alteração dos valores de tráfego e de QoS transportados nas mensagens PATH e RESV. Porém, a tecnologia ATM até a UNI 4.0 não suporta essa característica, ou seja, caso seja necessário um novo valor de QoS, será preciso que se estabeleça um novo canal virtual para suportar a nova requisição.

O processo de estabelecimento de canais virtuais ATM envolve operações complexas que demandam um tempo substancial. Além disso, o antigo canal virtual precisa ser fechado causando problemas com a abertura do novo dependendo da ordem em que as operações são realizadas. Se um novo canal for estabelecido antes que o antigo seja fechado pode ocasionar o problema de falsa falta de recursos, neste caso pode existir banda disponível para a nova requisição de QoS desde que fosse liberada pelo canal

exclusivos para aquelas mensagens. Se tomada a decisão pelo envio em conjunto com os canais de dados com QoS, apesar de não aumentar o número de conexões em uma sessão, alguns problemas podem surgir em situações de congestionamento devido ao descarte de pacotes. Uma variação desta solução, ainda em estudo, consiste em usar o canal de melhor esforço da sessão RSVP, porém com uma modificação no algoritmo de escalonamento para dar prioridade à sinalização RSVP em relação aos dados.

Com o uso exclusivo de canais virtuais para a sinalização RSVP não existe risco de congestionamento, pois podem ser feitas reservas de recursos para esse tipo de tráfego. O canal exclusivo para sinalização pode ser utilizado para mensagens de controle por sessão ou de forma agregada onde várias sessões compartilham uma conexão de controle. O primeiro esquema abre um canal virtual para as mensagens de controle de cada sessão; desta forma, aumenta muito o gasto de conexões e a quantidade de informações de controle por sessão. Com a agregação, as reservas de várias sessões são combinadas em um único canal com a perda de isolamento das reservas de tráfego, porém esta solução ainda é objeto de pesquisa [20].

3.5 Caracterização de tráfego e parâmetros de QoS

O segundo problema de adaptação RSVP sobre ATM envolve a integração da sinalização das duas arquiteturas [23]. Este problema pode ser dividido em duas partes: (1) a caracterização de tráfego; e (2) a especificação dos parâmetros de QoS; respectivamente os campos TSpec e RSpec de serviços do modelo IIS.

Tanto o modelo IIS quanto a tecnologia ATM implementam uma arquitetura de serviços. Essas arquiteturas possuem suas próprias definições de classes e parâmetros para caracterizar tráfego e especificar QoS. A solução para integração de sinalização consiste em traduzir as especificações IIS para os mecanismos de abertura de canais virtuais ATM.

Os serviços definidos pelo modelo IIS são traduzidos pelas categorias de serviço ATM [16]. Algumas categorias são mais apropriadas às restrições de um serviço IIS em particular, podendo oferecer um serviço mais otimizado.

armazenamento do balde b . Com esses valores, a rede ATM oferece um serviço com taxa nominal PCR e com tolerância à *jitter* (CDVT) determinada pela capacidade do balde.

Usando o serviço VBR(RT), duas taxas de serviço devem ser especificadas, máxima (PCR) e média (SCR), para modelar a característica de rajada do tráfego. Essas taxas são definidas pela taxa de pico p e pela taxa de geração r do *token-bucket* TSpec, respectivamente. O valor de capacidade do balde b especifica o tamanho máximo de rajada (MBS). O valor de tolerância à *jitter* (CDVT) é especificado pela relação entre a máxima rajada b e a taxa de serviço R garantida.

As seguintes relações podem ser derivadas das regras acima de forma que, o limite de atraso b/R e o tamanho das filas não ultrapassem o que foi especificado (onde θ representa a taxa máxima de transmissão) [11]:

$$R \leq PCR \leq \min(p, \theta)$$

$$r \leq SCR \leq PCR$$

$$b \leq MBS$$

O tráfego enviado que exceder os limites definidos na admissão (tráfego não conforme) deve ser marcado pelo policiamento com o parâmetro CLP (*Cell Loss Priority*) de forma a perder a prioridade e ser tratado como melhor esforço.

3.5.2 Especificações para o Serviço de Carga Controlada

Há duas opções para a implementação do serviço de carga controlada sobre ATM: (1) ABR; ou (2) VBR(NRT). Os serviços CBR e VBR(RT) seriam utilizados com potencial desperdício de recursos. A categoria UBR não oferece garantias mínimas para os requisitos do serviço de carga controlada.

Com a utilização da classe ABR, pode ser definido um valor positivo para a taxa mínima de envio (MCR) para garantir as porcentagens definidas pelo serviço mesmo em condições de congestionamento. Para diminuir o descarte da maioria dos pacotes, o parâmetro razão de perda (CLR) deve ser especificado com valor pequeno.

arquiteturas de serviço e foram abordados os principais problemas que surgem devido as diferenças entre os paradigmas. Foram apresentadas soluções para o gerenciamento de conexões e o suporte ao controle de QoS com o objetivo de amenizar as diferenças das arquiteturas. Contudo, conclui-se que as soluções sobrepostas apresentam deficiências e que as principais questões estão sendo abordadas em pesquisas realizadas atualmente.

Algumas soluções sobrepostas (p. ex. IP Clássico) só permitem que máquinas se conectem diretamente com outras via ATM caso ambas pertençam a uma mesma rede lógica. Caso contrário, elas devem se comunicar através do roteamento tradicional, mesmo que possam se conectar diretamente através da rede ATM. Desta forma, o roteador acaba se tornando um provável ponto de contenção e falha, não se aproveitando adequadamente a infra-estrutura de enlace.

O envio de dados *multicast* para ser eficiente deve considerar a real topologia da rede [29]. Porém, a abstração em nuvem esconde a topologia, gerando caminhos mais longos e até mesmo duplicação de tráfego. Nas soluções sobrepostas, o *multicast* é implementado através de servidores que indicam os pontos de saída da rede ATM para os quais devem ser estabelecidos canais virtuais [16]. Os canais virtuais são estabelecidos com a complexa sinalização ATM que introduz atrasos na inicialização da conexão.

Ainda que algumas soluções sobrepostas permitam suportar QoS, nenhuma dessas soluções é satisfatória. Em seu projeto inicial, essas soluções estabeleciam apenas conexões do tipo melhor esforço e não suportavam meios para informar parâmetros de QoS para o enlace ATM. Atualmente, essas soluções suportam QoS, mas sofrem do problema de heterogeneidade.

A natureza orientada por conexão da tecnologia ATM dificulta o suporte das características de tráfego que estão surgindo com demanda de QoS dinâmica. Tendências apontam para a modificação do padrão ATM para suportar conexões mais flexíveis que serão controladas pelas camadas superiores (rede/transporte OSI) com agregação de tráfego de vários canais virtuais através da utilização de filas individuais para escalonamento. Esses assuntos serão abordados nos próximos capítulos.

Capítulo 4

Modelo de Comutação por Rótulos

4.1 Introdução

Esse capítulo apresenta o modelo de comutação por rótulos que utiliza informações de roteamento da camada de rede OSI para o encaminhamento dos dados na camada de enlace OSI. Esse novo modelo de comutação surgiu como uma possível solução para aprimorar a integração da arquitetura IP com a tecnologia ATM. Seus mecanismos de operação foram projetados com o objetivo de evitar alguns dos problemas característicos dos modelos sobrepostos apresentados no capítulo anterior.

A comutação por rótulos se originou de pesquisas cujos objetivos eram de oferecer toda potencialidade ATM sem adicionar a complexidade de protocolos e de servidores utilizadas em outras soluções (ver apêndice A). Entre seus objetivos estão a padronização de um protocolo de distribuição de rótulos para fluxos *unicast/multicast*, definição de rotas explícitas para gerenciamento de tráfego e hierarquia de comutação.

A seção 4.2 apresenta a arquitetura de comutação por rótulos e discute seus mecanismos e características. Nas três subseções subsequentes são apresentados os principais tipos e implementações da comutação por rótulos. Alguns mecanismos desta nova solução estão em fase de padronização, especialmente aqueles relativos ao controle de QoS que ainda não estão estáveis. O suporte ao modelo IIS e modelos de QoS são abordados na seção 4.3 em seu estágio de desenvolvimento.

O rótulo de comutação é utilizado apenas para a associação de fluxos de redes à comutação no nível do enlace, isto é, sua semântica não representa uma codificação direta de um endereço de rede. O rótulo tem significado local para cada enlace; deste modo, as associações de rótulos com fluxos podem ser diferentes em cada elemento de rede.

4.2.1 Mecanismos de operação

Uma rede de comutadores por rótulos pode ser definida como um conjunto de comutadores com capacidade de processamento no nível de rede, denominados de elementos de rede, ou simplesmente, nó. Como mostra a Figura 4.1, um caminho de transmissão é uma seqüência de nós (1 – 4) por onde os pacotes de um fluxo atravessam a rede de comutação; o nó por onde são recebidos esses pacotes é denominado de nó de entrada (nó 1); o nó que envia os pacotes para uma máquina não pertencente à rede de comutação por rótulos é denominado de nó de saída (nó 4). Um nó considera como vizinho superior (*upstream*) o seu adjacente que pertence ao caminho de transmissão mais próximo ao nó de entrada, e como vizinho inferior (*downstream*) o seu adjacente do caminho mais próximo ao nó de saída. Na figura, o nó 3 tem como superior o nó 2, e como inferior o nó 4.

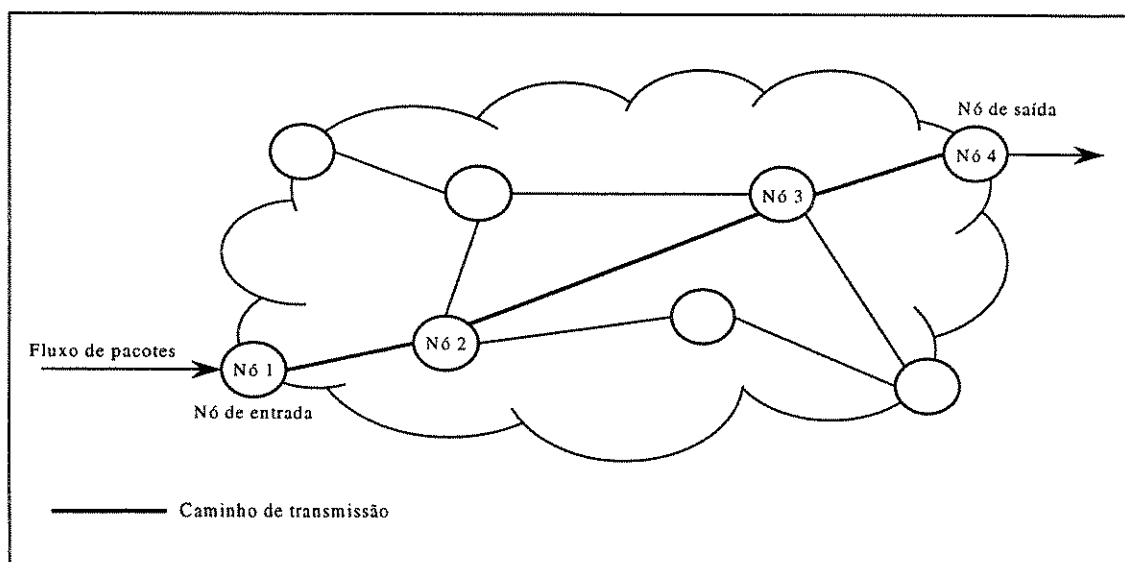


Figura 4.1: Rede de comutadores por rótulos de roteamento

caminhos comutados para gerenciamento de tráfego entre rotas; ou (2) de forma automática utilizando um mecanismo para distribuição de rótulos.

No modo de definição automático, o mecanismo de distribuição de rótulos é implementado por um protocolo de distribuição de rótulos – LDP (*Label Distribution Protocol*). Com esse protocolo, os nós enviam mensagens contendo o rótulo (um inteiro fixo) e os campos de cabeçalho de um pacote como identificação da FEC.

Por especificação da arquitetura, os caminhos comutados são válidos apenas dentro do domínio da rede de comutação [26]. Se um pacote atravessar mais de um domínio de comutadores por rótulos, entre esses domínios deve haver um novo processamento na camada de rede. Essa limitação pode ser contornada com o conceito de hierarquia de rótulos. Um conjunto de rótulos associado aos domínios atravessados por um fluxo é definido e armazenado em forma de pilha. Conforme as células ou quadros sejam transmitidas pelos domínios, os rótulos são desempilhados e usados para comutação dentro do respectivo domínio.

A comutação por rótulos engloba o conceito de agregação. Um rótulo pode agrupar uma quantidade variável de fluxos sendo transportados por um único caminho comutado. A agregação pode ser realizada por prefixos de endereços IP, por rotas de destino (nós de saída), por portas (TCP/UDP) de aplicativos, entre outras formas. A escolha do nível de agregação pode representar uma economia dos recursos de comutação, como por exemplo, o número de conexões utilizadas.

No estabelecimento de um caminho comutado as conexões entre cada nó são estabelecidas de forma independente com valores de QoS próprios. Desta forma, um caminho comutado pode possuir ramos com QoS variada, resolvendo os problemas de heterogeneidade e comportamento dinâmico dos modelos sobrepostos.

4.2.2 Comutação por rótulos em ATM

Antes de ser generalizado para qualquer meio de enlace, o conceito de comutação por rótulos foi desenvolvido para tornar mais eficaz as implementações dos protocolos de redes na arquitetura ATM.

A primeira solução não é eficiente pois a atribuição de um caminho virtual a um agregado pode desperdiçar muitos canais virtuais. Além disso, essa solução representa apenas uma forma de organizar o tráfego nos campos VPI/VCI e não uma agregação real. A agregação com VP requer o uso de um protocolo distribuído para distinção de canais virtuais de modo que dois fluxos provenientes de fontes distintas não se sobreponham.

A segunda solução desperdiça parte da largura de banda pois utiliza alguns bits da célula ATM para multiplexar os canais. Com isso, apenas um pequeno número de conexões podem ser multiplexadas. Além disso, essa camada não é utilizada como padrão para envio de dados em redes de computadores.

A camada de adaptação AAL5 é amplamente utilizada para transmissão de dados devido ao seu método econômico de encapsulamento. Porém, a AAL5 não é capaz de multiplexar em um único destino o tráfego proveniente de vários outros canais. Essa característica causa o problema de intercalação de células: se células de diferentes fluxos forem misturadas em um mesmo canal virtual não há forma de separação. Desta maneira, os rótulos não poderiam ser associados a fluxos agregados. Como visto na parte de comutação VC da figura, é necessário um mecanismo para evitar a intercalação das células com VPI 6 pois a camada AAL5 não é capaz de multiplexar o tráfego em um único canal virtual.

A terceira solução implementa uma modificação no envio das células de forma a habilitar a camada AAL5 multiplexar canais virtuais. As células que formam um quadro AAL5 só são transferidas para o canal virtual de saída na forma de um bloco sequencial e compacto após a chegada da última célula que forma o quadro e que é marcada com o bit de término AAL5. As primeiras células esperam a chegada das outras na memória da porta de entrada ou de saída. Apesar de evitar o problema de intercalação de células, esta solução ainda é objeto de pesquisas pois suas desvantagens (como o bloqueio de tráfego, aumento do uso de memória e rajadas de dados) devem ser superadas [27].

No paradigma de associação de rótulos com fluxos de rede, as informações de estado para comutação das células ATM são armazenadas de forma temporária e local em cada nó do caminho de transmissão. Essas informações podem ser atualizadas através de

deve ser estabelecida uma conexão. Com isso, canais virtuais podem ser abertos até mesmo para rotas ociosas. Além disso, mudanças nas tabelas de roteamento podem levar a problemas de estabilidade e restabelecimento de conexões devido aos algoritmos utilizados para convergência de rotas.

Nas subseções abaixo são apresentados dois exemplos de implementação, uma orientada por tráfego e outra por topologia, as arquiteturas *IP Switch* e *Tag Switch*, respectivamente. Essas implementações são os principais representantes de seus tipos e apresentam diferentes características de aproveitamento da tecnologia ATM.

4.2.4 Arquitetura *IP Switch*

A arquitetura *IP Switch* [28,29] é composta de um comutador ATM padrão⁴, sem a sinalização de controle acima da camada AAL5 e adicionado em seu lugar um módulo de controle IP. Deste modo, são removidos os protocolos de roteamento e o gerenciamento de conexões (protocolos como o PNNI e a sinalização Q.2931), cujas funções são realizadas no nível de rede pelos protocolos IP, possibilitando uma maior flexibilidade e robustez. Com isso, consegue-se a eliminação da redundância de funcionalidades entre as camadas de rede e enlace.

4.2.4.1 Características Gerais

O módulo de controle IP é auxiliado por dois protocolos de gerenciamento de conexões e possui um mecanismo de classificação de fluxos (Figura 4.3). O protocolo IFMP [30] (*Ipsilon Flow Management Protocol*) controla o estabelecimento de canais virtuais entre comutadores adjacentes, substituindo o gerenciamento orientado por conexões ATM. O protocolo GSMP [31, 32] (*General Switch Management Protocol*) possibilita o controle dos circuitos de comutação diretamente na malha ATM. O mecanismo de classificação escolhe com critérios locais quais fluxos devem ser

⁴ Qualquer outro tipo de enlace pode ser utilizado, porém, a literatura considera apenas o ATM

Se o fluxo ao qual pertence o pacote analisado for escolhido pelo classificador para ser comutado, o módulo de controle executa diversos procedimentos (através de mensagens GSMP e IFMP) para configurar o comutador de modo que os próximos pacotes daquele fluxo não necessitem mais do tratamento no nível da camada de rede. Caso contrário, os pacotes dos fluxos não escolhidos continuam sendo recebidos e enviados, via porta de controle, para tratamento pelo nível de rede.

Quando um fluxo obtém uma conexão própria para comutação, o *IP Switch* armazena as informações de cabeçalho daquele fluxo em suas variáveis de estado, e passa a transmitir seus dados diretamente por essa conexão. Na transmissão não são utilizadas essas variáveis armazenadas ou qualquer outra forma de encapsulamento, ou seja, os dados são transportados diretamente pela AAL5. No final de um caminho comutado, o último *IP Switch* remonta os pacotes com os dados que chegam na conexão e adiciona as informações armazenadas nas variáveis de estado.

O *IP Switch* suporta *multicast* sem qualquer modificação nos protocolos utilizados atualmente na arquitetura IP (inclusive o serviço multiponto-multiponto), ao contrário do padrão ATM que necessita de protocolos próprios. Isto porque, o protocolo GSMP modifica as tabelas de comutação para estabelecer cada ramo da árvore de difusão de forma individual sem a sinalização ATM. Os ramos podem ser modificados a qualquer momento, inclusive quanto aos valores de QoS atribuídos aos ramos.

4.2.4.2 Protocolo de gerenciamento de fluxo (IFMP)

O protocolo IFMP tem a função de distribuição de rótulos na arquitetura *IP Switch*. O IFMP permite ao controle IP solicitar que um vizinho envie dados por um canal específico para comutação direta de um fluxo classificado. Além disso, esse protocolo permite que um elemento de rede descubra as máquinas de sua vizinhança e mantenha canais virtuais sincronizados com esses vizinhos depois de mudanças ocorridas no enlace.

No momento em que um fluxo é classificado (Figura 4.4 – A), esse fluxo é associado a um rótulo que identifica em qual canal virtual ele será enviado. Esse rótulo é

O tipo de fluxo define o grau de granulação do tráfego associado a um canal em particular. Cada tipo especifica quais campos do cabeçalho IP serão usados como método de classificação, os quais são armazenados em variáveis de estado do controle IP associadas a um canal virtual. Todos os pacotes cujos campos de cabeçalho sejam idênticos pertencem ao mesmo fluxo. Vários tipos de fluxos podem ser definidos no IFMP, porém apenas dois tipos estão definidos: FT1, pacotes em que os valores dos endereços IP e das portas TCP/UDP de entrada/saída sejam idênticos, ou seja um fluxo entre dois processos; e FT2, pacotes em que somente os valores de endereços IP sejam idênticos, ou seja, um fluxo entre duas máquinas.

Na operação de armazenamento das variáveis de estado associadas a um canal virtual, o campo TTL é também armazenado para os dois tipos de fluxos subtraído de um, isto com o objetivo de prevenir rotas cíclicas, já que quando comutado não há modo de verificação deste campo. O último *IP Switch* remonta o pacote com as variáveis de estado e o campo TTL atualizado. Esse armazenamento também garante segurança já que não permite que um fluxo seja utilizado para enviar pacotes não acordados previamente.

4.2.4.3 Protocolo geral de gerenciamento (GSMP)

O protocolo GSMP permite ao controle da arquitetura *IP Switch* estabelecer e gerenciar os circuitos na malha de comutação ATM. Entre outras funções, o GSMP permite a obtenção de estatísticas de uso do canal de comunicação, a requisição de parâmetros de configuração do comutador e a notificação sobre ocorrência de eventos assíncronos, tais como erros relacionados a alguma porta ou conexão. O controle da malha de comutação e as requisições de informação são enviadas através de mensagens próprias do protocolo.

O envio de mensagens de gerência permite ao controle IP estabelecer e remover circuitos que conectam a malha de comutação. Cada mensagem conecta um canal de entrada com um de saída, ou seja, cria um ramo e faz associação individual de prioridade. As mensagens GSMP permitem tratar a adição e remoção de ramos de forma simplificada, sem a complexidade da sinalização ATM. Para estabelecer conexões *multicast*, entretanto, são necessárias tantas mensagens GSMP quanto o número de ramos

4.2.5 Arquitetura *Tag Switch*

A arquitetura *Tag Switch* [35] propõe um esquema de comutação onde os caminhos virtuais são construídos de forma antecipada, isto é, quando um fluxo de dados começa a ser transmitido seus pacotes já podem ser comutados. A construção antecipada é realizada de forma integrada com o mecanismo de roteamento que utiliza suas tabelas de rotas para associação com rótulos. Desta forma, a arquitetura tem a característica de ser orientada à topologia.

4.2.5.1 Componentes da Arquitetura

Os elementos que formam a arquitetura são apresentados a seguir juntamente com suas funções:

Roteadores de borda: delimitam a área abrangida pela rede por rótulos e fazem vizinhança com redes de outra tecnologia. Têm a função de executar a avaliação dos cabeçalhos dos pacotes no nível de rede e atribuir-lhes um rótulo.

Comutadores por rótulos: máquinas pertencentes ao meio da rede. Executam o repasse de células através de seus rótulos. Trocam informações sobre rótulos entre si e com os roteadores para a construção dos caminhos comutados.

Protocolo de distribuição de rótulos (TDP – *Tag Distribution Protocol*): distribui para os elementos acima citados as informações sobre a associação de rótulos com rotas na rede.

A operação básica em uma rede *Tag Switch* inicia-se com os roteadores de borda obtendo informações sobre as rotas usadas nas tabelas de roteamento da rede. Os roteadores e os comutadores trocam essas informações para construção dos caminhos comutados. Para cada entrada existente na tabela de rotas é associado um rótulo. Esse rótulo é distribuído na rede via o protocolo de distribuição de rótulos. Um roteador de borda armazena essas associações para que possa atribuir o rótulo correto a um pacote de dados pela análise de seu cabeçalho e identificando sua rota. O pacote é então transmitido

4.2.5.2 Associação de rótulos

Uma das funções do componente de controle é a criação de rótulos baseada no endereço de destino dos pacotes, cujas rotas são formadas antecipadamente pelos protocolos de roteamento. Há três formas de associação de rótulos para o preenchimento da TIB na arquitetura *Tag Switch*. Todas as formas associam as entradas da FIB a um rótulo e, de acordo com a sua forma, seguem os seguintes passos (Figura 4.5):

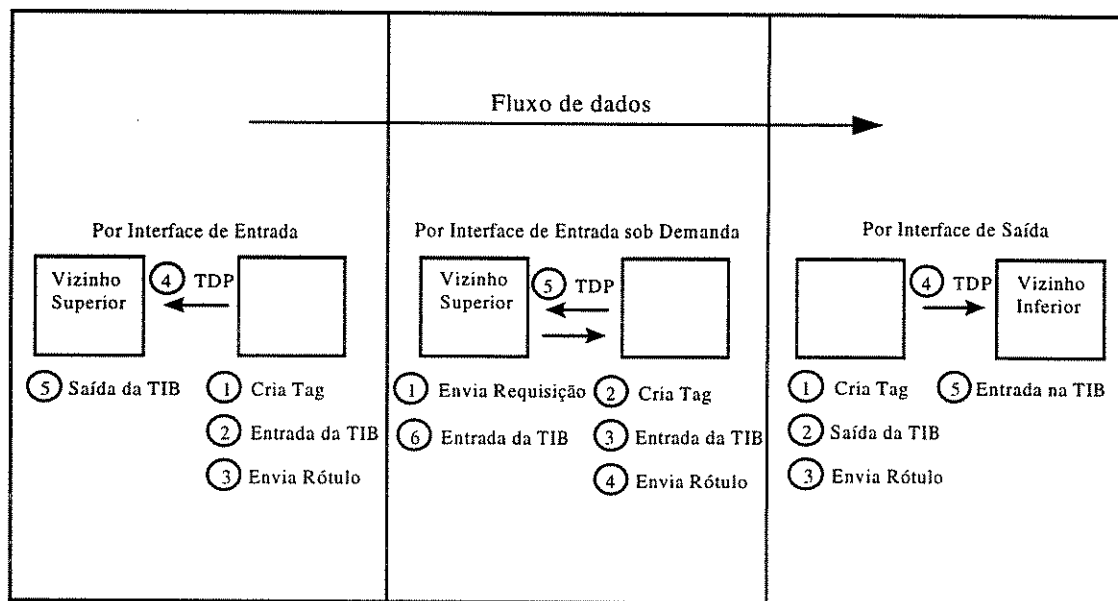


Figura 4.5: Formas de Preenchimento da TIB

- Associação pela interface de entrada: o rótulo criado é armazenado como rótulo de entrada na TIB, e essa associação é enviada pelo protocolo de distribuição aos vizinhos superiores. Quando um vizinho recebe essa associação, os campos da tripla rótulo de saída e interface de sua TIB são preenchidos com o rótulo enviado e a porta de saída.
- Associação sob demanda pela interface de entrada: ao invés de enviar o rótulo é enviada uma requisição de associação. Assim o vizinho superior cria um rótulo associado à rota, preenche sua TIB com o valor do rótulo de saída, e envia o resultado. Recebida a resposta, a TIB e a FIB locais são preenchidas com o valor do rótulo de entrada.

no preenchimento da TIB para construção dos caminhos comutados. O protocolo de distribuição de rótulos (TDP) é utilizado pelas três formas de associação de rótulos usadas no preenchimento da TIB, sendo utilizado por todos os módulos do componente de controle.

O protocolo de distribuição deve garantir a entrega das mensagens por motivo de consistência das tabelas de comutação. Para isso, o protocolo abre sessões para o envio das mensagens cujo recebimento deve ser confirmado. As mensagens do protocolo devem ser enviadas pelo nível de transporte através de um protocolo confiável para que se obtenha garantias de entrega sequencial. A codificação dos rótulos e das rotas é enviada em campos reservados do protocolo juntamente com o código de operação desejada.

O protocolo de distribuição de rótulos também fornece meios para notificação, requisição e liberação das associações de rótulos com rotas de acordo com a dinâmica das tabelas de roteamento. Além disso, oferece operações de suporte, gerenciamento e notificação de erros da própria sessão de transmissão do protocolo.

4.3 Comutação por rótulos e o modelo IIS

A comutação por rótulos resolve parte dos problemas apresentados na adaptação dos modelos clássicos e IIS sobre ATM padrão, como a heterogeneidade e o comportamento dinâmico de QoS. Porém, o suporte aos estilos RSVP ainda necessitam de aprimoramentos. Soluções para os problemas relacionados com QoS e agregação ainda não foram padronizadas pelo grupo MPLS/IETF. Características de algumas arquiteturas facilitam a adequação aos mecanismos do modelo IIS; porém, são implementados de forma complexa gerando problemas de compatibilidade.

Os trabalhos existentes se dividem em duas áreas principais: (1) suporte ao protocolo RSVP; (2) suporte para classes de serviço e QoS. Os mecanismos de suporte RSVP têm o objetivo de estabelecer os caminhos comutados para o tráfego dos fluxos prioritários de acordo com as mensagens de controle. O suporte às classes de serviço se baseiam em modelos de QoS do protocolo GSMP.

para as mensagens de controle ICMP/IGMP, o IETF vem desenvolvendo mecanismos complexos de tratamento do campo TTL e extensões ao protocolo ICMP para MPLS⁵.

Em uma arquitetura orientada por tráfego, o caminho comutado é construído junto com o trânsito dos primeiros pacotes de um fluxo pela rede. Desta forma, esses pacotes podem armazenar a quantidade de máquinas que atravessaram em um caminho comutado e decrementar o campo TTL do fluxo, como é tradicionalmente realizado pelo protocolo ICMP.

Para que um fluxo comece a receber tratamento diferenciado é necessária a separação do fluxo em um caminho comutado estabelecido pela sessão RSVP. Para isso, os fluxos RSVP devem ser atribuídos a rótulos de acordo com o estilo utilizado. A associação de rótulos com fluxos deve ser realizada pelo receptor (por interface de entrada). Se outro método fosse escolhido (as fontes alocando rótulos), poderia haver condições de concorrência pelos rótulos disponíveis pois a alocação pelo receptor é o método já padronizado em conexões melhor esforço [36].

O estabelecimento de conexões *IntServ* é realizado através de mensagens RESV. Para isso, foi criado o objeto RSVP_LABEL para distribuir a associação entre rótulos e fluxos RSVP. A associação é determinada pelo primeiro receptor das mensagens PATH associadas que então a distribui aos outros receptores através do uso de mensagens PATH. O rótulo encapsulado no objeto RSVP_LABEL é associado ao fluxo diretamente pelos campos FILTER_SPEC (mensagem RESV) e SENDER_TEMPLATE (na mensagem PATH). Caso nenhum rótulo seja recebido para um determinado fluxo, este será enviado pelo caminho comutado melhor esforço. A proposta desse objeto se baseia na arquitetura *Tag Switch* e está sendo padronizada para a comutação por rótulos.

A associação de rótulos no estilo FF é simples, cada fluxo necessita de uma conexão exclusiva implicando em um rótulo para cada fonte selecionada. Nos estilos compartilhados (SE ou WF), a associação deve explorar a economia de rótulos, pois esse

⁵ Até o período de redação deste trabalho não existam documentos padronizados para extensão ICMP (veja site MPLS/IETF <http://www.ietf.org/mpls>).

podendo configurar um nível de serviço e QoS. Desta maneira, o GSMP permite configurar de forma razoável o conceito de classes de serviço no comutador.

O sistema controle/comutador deve estar apto a satisfazer os requisitos de QoS da tecnologia de rede. Para isso, as funções de gerenciamento de recursos devem estar divididas de alguma forma entre o controle e o comutador. Além disso, precisa ser decidido se o modelo de serviço deve ser genérico ou possuir definições específicas dos parâmetros de QoS.

Um gerenciamento de serviços genérico é potencialmente complexo, pois há necessidade de abstração dos mecanismos de controle de QoS implementados no comutador. Esse tipo de gerenciamento é utilizado nas versões 1.1 e 2 do protocolo GSMP. O gerenciamento específico utiliza serviços padrões de modelos como IIS, DS, ATM Fórum, entre outros. No estabelecimento de conexões, o controle envia a especificação do serviço para que o comutador ajuste seus mecanismos de controle (reguladores, filas de escalonamento, etc.).

Se o gerenciamento for localizado apenas no comutador, o nível de rede terá dificuldades de controlar recursos compartilhados, aumentando a complexidade de tarefas como o balanceamento de tráfego. Além disso, esse modo se afasta dos modelos de serviço IP onde o controle é feito pela camada de rede (caso *IntServ*). Por outro lado, se o gerenciamento for localizado apenas no controle de rede, algumas tarefas de comutação (classificação e escalonamento) podem se tornar inconsistentes devido as particularidades do *hardware* do comutador. Uma solução mista em que o gerenciamento de recursos é dividida entre o comutador e o controle, que utiliza um modelo de serviço específico, será apresentado no próximo capítulo.

4.4 Conclusões

As soluções sobrepostas apresentam problemas que dificultam a adaptação otimizada dos protocolos IP sobre a tecnologia ATM. Nessas soluções, as dificuldades para oferecer QoS, a complexidade dos mecanismos de controle e a duplicidade no envio

Capítulo 5

Comutação com Serviços Integrados

5.1 Introdução

A comutação por rótulos foi desenvolvida para facilitar a integração IP/ATM, como visto no capítulo anterior. Contudo, existem alguns aspectos que devem ser considerados: falta de tratamento quantitativo do tráfego sensível a tempo, problemas de suporte à agregação de fluxos, utilização excessiva de identificadores de conexões (ocasionando escassez de rótulos) e o reconhecimento de mensagens de controle.

A arquitetura de Comutação com Serviços Integrados – Comutação IS – é proposta com objetivo de oferecer alternativas para solucionar esses problemas. Para isso, exerce um controle efetivo sobre QoS que integra o modelo IIS e o protocolo RSVP com a capacidade de comutação direta dos fluxos no nível de enlace. Além disso, propõe novas políticas de agregação de tráfego para diminuir a quantidade de canais virtuais utilizados pelos fluxos comutados.

A arquitetura do novo modelo de comutação é proposta na seção 5.2 com base na arquitetura proprietária *IP Switch* cuja escolha é discutida nesta seção. A seção 5.3 analisa a política de gerência de conexões utilizada pela Comutação IS e resultados obtidos em simulação. A seção 5.4 mostra os mecanismos de suporte para o protocolo RSVP. A seção 5.5 esboça um modelo de escalonamento que integra os mecanismos do modelo IIS com a camada de enlace. A seção 5.6 introduz o conceito de comutação adaptativa e mostra os resultados obtidos em simulação da nova arquitetura.

5.2 Comutação por Rótulos com Serviços Integrados

A comutação por rótulos permite que os protocolos IP possam ser utilizados com a tecnologia ATM sem a necessidade de soluções complexas de adaptação [29]. Os mecanismos da comutação por rótulos agilizam o envio de dados, pois evitam os atrasos introduzidos pela camada de rede no tratamento dos cabeçalhos IP. Contudo, a solução de

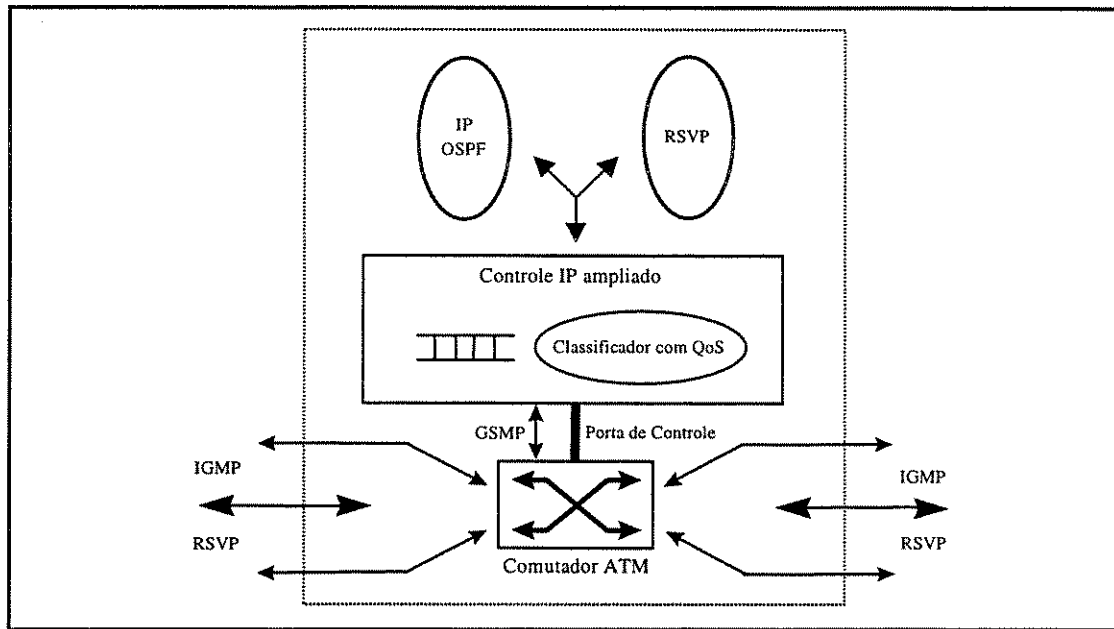


Figura 5.1: Arquitetura IS em um comutador por rótulos

O protocolo GSMP (versão 2 [32]) é utilizado para controlar o comutador ATM de modo a estabelecer conexões, controlar o escalonamento e suportar níveis diferenciados de QoS. O protocolo GSMP torna o controle IP independente do *hardware* e vice-versa. Desta forma a sinalização ATM pode ser substituída e a comutação ser diretamente controlada pela camada de rede. Conexões *multicast* podem ser estabelecidas com QoS diferenciada por ramo na malha de difusão.

A manutenção da análise dos cabeçalhos dos primeiros pacotes IP de um fluxo mantém a robustez e eficiência dos protocolos IP, pois mantém a filosofia original da arquitetura. Ainda assim, o estabelecimento de caminhos comutados por mensagens RSVP permite que o processo de envio de fluxos com QoS seja agilizado pois canais virtuais específicos podem ser associados para tratamento desses fluxos. Ao contrário, uma arquitetura orientada por topologia perde em flexibilidade, pois não analisa as informações contidas nos cabeçalhos IP.

Como demonstrado em simulação [44], arquiteturas de comutação por rótulos podem demandar uma quantidade considerável de identificadores de conexões. Esse recurso é escasso nas redes ATM pois a codificação é feita nos campos VPI/VCI do cabeçalho das células sendo limitada pelo tamanho máximo de 28 bits. Além disso, em

O mecanismo de controle da arquitetura de comutação IS introduz uma modificação em relação à arquitetura *IP Switch* para reconhecimento de fluxos com necessidade para controle de QoS. A identificação desses fluxos será realizado por comunicação através de objetos RSVP ao módulo de controle IP (como apresentado na próxima seção). Essa modificação permite a associação de cada fluxo com QoS a uma conexão individual de forma que um caminho comutado seja estabelecido para tratamento deste tráfego.

O tratamento especial de fluxos compromete o espaço de rótulos pois aumenta o uso de canais virtuais, aproximando-se dos resultados obtidos com o uso do tipo de fluxo FT1. A Comutação IS recorre a recursos como agregação de canais virtuais para dar suporte aos estilos RSVP e utiliza políticas orientadas por topologia para o estabelecimento de conexões.

Esta seção apresenta e confirma os resultados obtidos em simulação com a utilização da ferramenta *issim* [44] (ver apêndice B). A reprodução de resultados por simulação busca entender a influência da variação de parâmetros de comutação no desempenho da arquitetura para formar uma base de avaliação e posterior comparação do desempenho do comutador depois que mudanças no módulo de controle IP sejam introduzidas. Essas mudanças visam aumentar a agregação de tráfego para diminuir o número de canais virtuais utilizados como discutido nas próximas seções.

5.3.1 Estabelecimento de Conexões

A Comutação IS mantém a utilização do protocolo IFMP [30] (*Internet¹ Flow Management Protocol*) como mecanismo de distribuição de rótulos para fluxos sem QoS. Para fluxos com QoS tanto *unicast* como para *multicast*, o protocolo RSVP será utilizado para estabelecimento (veja próxima seção). Para isso, é proposto um novo tipo de mensagem IFMP/RSVP com conteúdo idêntico para distribuição de rótulos (veja próxima seção para definição do objeto RSVP).

¹ Este trabalho propõe a utilização da palavra *Internet* ao invés do nome proprietário.

A diminuição da quantidade de rótulos pode ser atingida com a alteração dos parâmetros internos à comutação. Porém, isso eleva a quantidade de eventos (abertura e fechamento) para gerenciamento de conexões. Esta seção mostra os procedimentos de obtenção e análise dos resultados da variação de parâmetros de comutação em simulação, e como essas variações alteram a quantidade de canais virtuais e o número de eventos por segundo no comutador.

A redução na quantidade de rótulos utilizados pode ser realizada pela alteração dos parâmetros X e Y (aumento de X ou redução de Y) e, principalmente, com a diminuição do intervalo de tempo permitido entre pacotes de um fluxo sem causar a desativação da conexão (FTO – *flow timeout*). Essa redução de rótulos acarreta um aumento proporcional do número de eventos internos ao comutador (abertura e fechamento de conexões) que a maioria dos comutadores não consegue suportar.

Os parâmetros utilizados neste trabalho são escolhidos de modo a maximizar a necessidade por recursos considerados escassos em um comutador real, como por exemplo, a quantidade de canais virtuais e a carga de gerenciamento (abertura e fechamento) destes canais. Essa maximização força a máquina de comutação operar próxima ao seu limite e torna mais claro os possíveis efeitos causados por alterações que buscam a melhoria no desempenho do sistema.

Um fluxo comutado é caracterizado pela definição do seu tipo (FT2, neste caso), um valor de *timeout* (FTO) para desativar conexões ociosas e os critérios de classificação X/Y. Tomando-se como base a análise já realizada nos artigos de simulação da literatura e visando o modo de operação limite, foram escolhidos os seguintes valores iniciais para esses parâmetros: $X/Y = 5/40$ e $FTO = 90s$.

Com esses valores de operação foram obtidos os resultados apresentados na Tabela 5.1 para cada arquivo de tráfego usado como entrada no simulador. A média de eventos que são gerados no comutador e o número de canais usados foram calculados por normalização linear para uma taxa de 1Gbps através do uso dos seguintes valores obtidos por simulação: número médio de conexões/desconexões, média de conexões ativas e da taxa de transmissão em cada arquivo de tráfego.

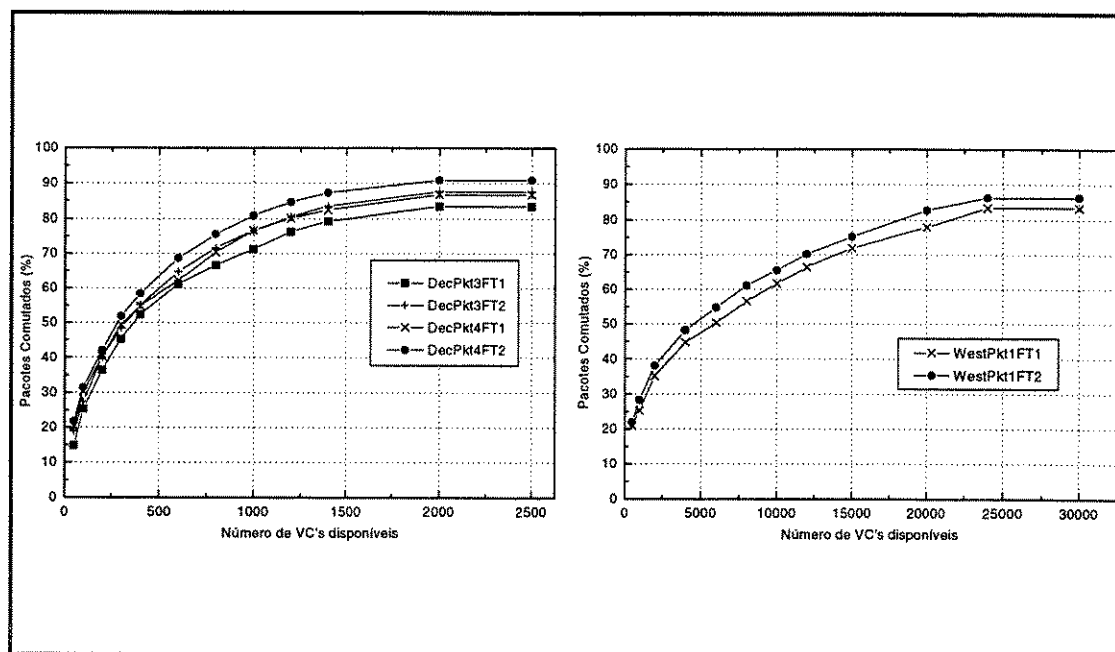


Figura 5.2: Porcentagem de comutação $\times$ VCs disponíveis ($X=5$ $Y=40$ $FTO=90$)

O aumento do número de conexões disponíveis para comutação proporciona um aumento no número de fluxos que podem ser acomodados, e desta forma, aumenta a porcentagem de pacotes comutados. Por outro lado, a porcentagem de comutação diminui com a redução de canais ativos no comutador. A porcentagem de comutação aumenta até atingir um valor máximo associado a uma quantidade máxima de conexões usadas por um dado padrão de tráfego. Essa quantidade máxima de conexões deve ser maior que o número de canais virtuais disponíveis no comutador; caso contrário a comutação de alguns fluxos pode ser rejeitada, diminuindo a quantidade de pacotes comutados.

Os comportamentos observados nos dois tipos de padrões de tráfego são similares, a menos da variação (em torno de uma ordem de grandeza) que ocorre no número de VC's usados devido à diferença nas taxas de tráfego. Como o objetivo das simulações visa a análise do uso do espaço de conexões do comutador, no restante dos gráficos apenas o arquivo *FixWest* é analisado, pois é o que apresenta maior volume de tráfego, diminuindo a possibilidade de erros devido à normalização em menor escala.

A Tabela 5.2 analisa a influência da variação do tempo FTO e do período de tempo Y necessário para caracterizar um fluxo comutado (fixando-se o número de

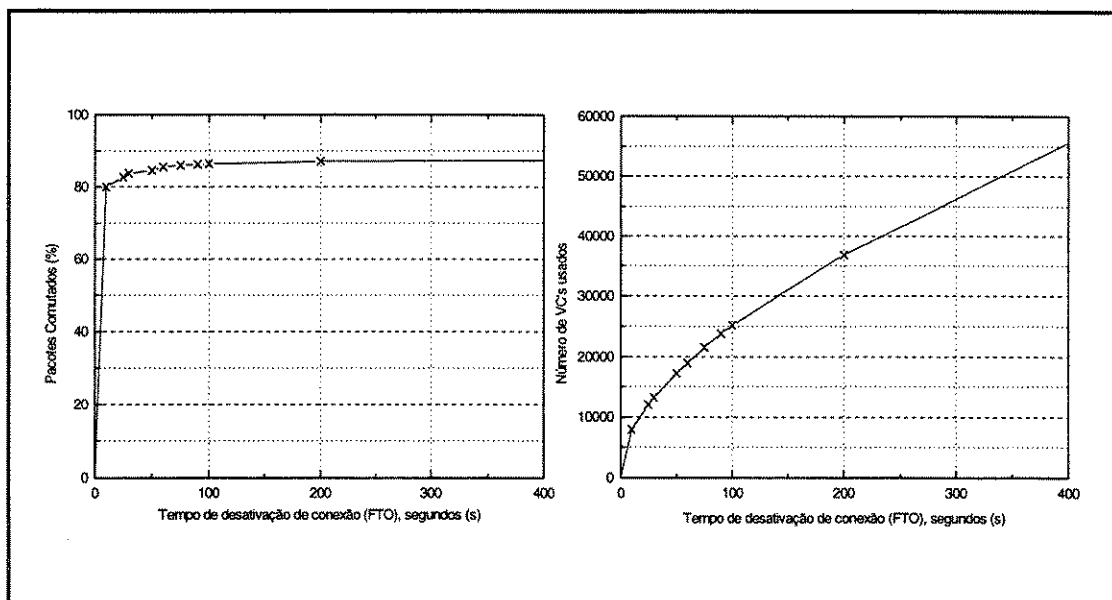


Figura 5.3: Efeito da variação do tempo de desativação de conexões (FTO)

Na Figura 5.4 são apresentados os gráficos de estabelecimento e desativação de conexões por segundo para os valores de FTO em 10 e 90 segundos. No começo da simulação só existem estabelecimentos de conexões, as desativações só começam a acontecer a partir do tempo de simulação superior ao FTO. Depois da inicialização, com a simulação já estabilizada, as curvas de ativação e desativação tendem a se manter próximas (na figura as duas curvas se sobrepõem a partir de 40 e 120 segundos respectivamente) estabilizando a curva total de atualizações da tabela de comutação.

Como a operação de gerência da tabela de comutação é um fator restritivo no desempenho de um comutador, devem ser usados valores de FTO maiores com o objetivo de se limitar uma possível sobrecarga. Porém, um aumento de FTO aumentaria o número de conexões usados no comutador (como mostra a Figura 5.3).

Os gráficos e a tabela mostram que para diminuição de canais virtuais o parâmetro FTO deve ser diminuído. Porém, de forma proporcional a essa diminuição ocorre um aumento da frequência interna de eventos no comutador. Além disso, a diminuição de FTO reduz drasticamente a porcentagem de pacotes comutados. Se o valor de FTO for mantido no ponto de estabilidade da Figura 5.3, o valor de Y deve ser diminuído para

agregação para diminuição do uso de canais virtuais com impacto positivo no número de eventos e na porcentagem de comutação do tráfego, como discutido nas próximas seções.

5.4 Suporte ao protocolo RSVP

Como foi visto nos capítulos anteriores, o modelo de comutação por rótulos resolve parte dos problemas relacionados ao gerenciamento de conexões RSVP, como a heterogeneidade e o comportamento dinâmico de QoS. Essas soluções se devem ao modo como é realizado o estabelecimento de conexões naquele modelo. Diferente do que ocorre no ATM padrão, que estabelece todos os ramos de uma conexão em uma operação única com QoS homogênea e estática, a comutação por rótulos estabelece cada ramo individualmente com recursos de QoS diferenciados para cada interface.

Apesar de solucionar estes problemas, a comutação por rótulos necessita de aprimoramento para resolver problemas como agregação, próprios do modelo IIS em sessões RSVP, e que são agravados em um ambiente baseado em comutação, como nas redes ATM. Particularmente, dois problemas específicos precisam ser resolvidos: (1) o suporte às conexões de controle RSVP; e (2) como o ATM, as arquiteturas de comutação por rótulos não suportam o conceito de conexões multiponto-multiponto necessário aos estilos compartilhados RSVP (SE e WF).

Cabe ressaltar que o suporte *multicast* na Comutação IS não traz modificações em relação à arquitetura *IP Switch*. Desta forma, a dinâmica das sessões RSVP é realizada pela política de abertura de conexões como definida pela Comutação IS, que permite a agregação de fluxos e o estabelecimento de conexões como necessário aos estilos RSVP.

5.4.1 Mensagens de Controle

Em uma arquitetura de comutação por rótulos, as mensagens de controle (PATH e RESV) podem ser comutadas em conjunto com o tráfego já classificado de um fluxo. Portanto, podem não receber tratamento adequado no nível de rede nos nós intermediários de uma rota, como especificado pelo protocolo RSVP.

tamanho de cabeçalho são especificados pelas informações contidas no descritor de fluxo (*FILTER_SPEC*).

5.4.2 Estilos de Reserva

Um fluxo RSVP é associado a uma fonte que envia dados a um identificador de sessão. Uma sessão pode conter vários fluxos (originados de várias fontes) sendo enviados a um conjunto de receptores. A reserva é realizada de acordo com o estilo de reserva escolhido, podendo ser exclusiva (estilo FF) ou compartilhada (estilo WF ou SE). No estilo compartilhado, uma conexão deve ser estabelecida de modo a agregar vários fluxos de uma mesma sessão.

Para agregação, é necessário dispor de conexões multiponto-multiponto onde vários endereços fontes são associados a um único canal virtual para envio de dados com compartilhamento de banda. No *IP Switch* e no padrão ATM, isto não pode ser implementado pois cada par fonte/destino representa um fluxo diferente (pares de endereços distintos) e logo são comutados em canais diferentes; neste caso não há como multiplexar as células em um único canal virtual.

Os estilos compartilhados são normalmente suportados pela comutação IS devido ao suporte de agregação de tráfego de diferentes canais virtuais nos comutadores com filas nas interfaces de saída. Em conexões WF, a máscara de endereços é definida de forma a aceitar o recebimento de pacotes de qualquer fonte. Em conexões FF e SE, as fontes precisam ser explicitamente indicadas.

O estilo SE sofre outro problema de adaptação na comutação por rótulos. Se dois receptores requisitam tráfego de subconjuntos diferentes da lista de fontes e esses fluxos forem agregados em um único canal, não há como separar tráfego na bifurcação dos dois receptores pois haveria necessidade de remontar os pacotes. Desta forma, os pacotes devem ser remontados e analisados para encaminhamento correto.

ocorre o objeto é reenviado a quem fez a requisição através de uma mensagem RESVERR, caso contrário, o objeto é reenviado através de uma mensagem PATH como confirmação do estabelecimento da conexão. Qualquer outro problema com a conexão o objeto pode ser enviado junto com as mensagens RSVP específicas, como por exemplo, com a de liberação explícita de conexão PATHTEAR.

5.5 Escalonamento Integrado

A comutação IS propõe o tratamento de tráfego integrado entre as camadas de rede e de enlace com quatro categorias de serviço disponíveis (figura 5.6): (1) garantido; (2) carga controlada; (3) melhor esforço de controle; e (4) comutado. A política de escalonamento é dividida em duas partes: o algoritmo de escalonamento de pacotes é executado nas portas de saída do comutador e o controle de admissão executado no controle IP de acordo com políticas mais complexas.

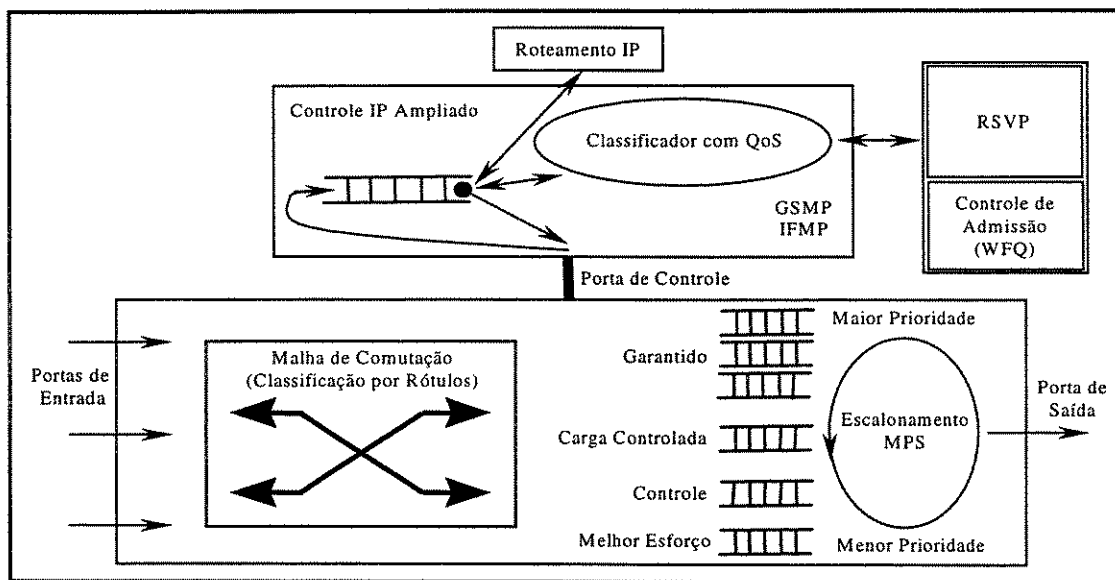


Figura 5.6: Escalonamento integrado nas camadas de rede e de enlace

Na malha de comutação o algoritmo escolhido para escalonamento foi o MPS (*Multi Priority Server*) [48] que compartilha a banda pela atribuição de níveis de prioridades. À categoria de serviço garantido são atribuídas as maiores prioridades para reservar recursos e garantir níveis de atraso [48]. Para a categoria de carga controlada é

Para implementar a agregação por topologia os parâmetros de classificação de fluxos devem ser ajustados de forma a se adaptar a mudança nos padrões de tráfego. Essas mudanças são muito comum na história da Internet e de certa forma esperada com a introdução de tráfego sensível a QoS.

Deste modo, o Comutador IS é proposto como um equipamento capaz de adaptar sua operação de acordo com o volume de tráfego sustentado e a necessidade por canais virtuais. O Comutador IS pode iniciar sua operação com um tipo de tráfego menos granulado como FT2, ou até mesmo FT1, identificando fluxos prioritários requisitados via RSVP.

Em determinado momento, baseado no limite de canais virtuais utilizados, o Comutador IS pode modificar a sua operação com o objetivo de agregar fluxos FT2 ou com maior grau de granulação. Para operar desta maneira a Comutação IS define novos tipos de fluxos e políticas de utilização de canais virtuais.

5.6.1 Agregação de Fluxos

Novos tipos de fluxos com agregação² – FT3 e FT4 – são propostos para diminuir o número de conexões utilizadas na comutação e permitir que os canais virtuais economizados sejam utilizadas para fornecer QoS sem sobrecarregar o sistema. O tipo de fluxo FT3 utiliza políticas de agregação, por domínios IP ou por prefixos de endereços, para associar um maior número de fluxos a um único canal virtual. O tipo de fluxo FT4 representa um grau maior de agregação, associando todo tráfego com um único destino (ponto de saída comum na rede de rótulos) a um único canal virtual. O FT4 aproxima-se mais da operação de comutadores orientados a topologia como o *Tag Switch*.

Como sugere a Figura 5.7, a definição na arquitetura *IP Switch* de uma política de classificação de fluxos que considere informações sobre topologia (FT3 e FT4) pode aproximar o desempenho dos dois tipos de arquiteturas de comutação. Com a variação

² A agregação de fluxos difere do conceito de agregação de tráfego para os estilos RSVP.

5.6.2 Roteamento sem Classes de Endereços

O modelo de roteamento sem classes de endereço – CIDR [47] (*Classless Inter Domain Routing*) – tem o objetivo de agregar em uma única rota endereços IP que tenham em comum uma parte da máscara de rede (contínua justificada à esquerda, como mostra a Figura 5.8). O CIDR engloba o conceito de endereçamento em classes A/B/C da arquitetura TCP/IP pois não diferencia as classes para realizar a operação de agregação.

Endereços Classe C:	0	8	16	24
192.32.1.0	11000000	00100000	00000001	00000000
192.32.2.0	11000000	00100000	00000010	00000000
Máscara de Rede Classe C:				
255.255.255.0	11111111	11111111	11111111	00000000
Endereço CIDR:	tamanho do prefixo →			
192.32.0.0/16	11000000	00100000	00000000	00000000

Figura 5.8: Endereçamento CIDR

Um endereço CIDR é formado por um prefixo do endereço IP mais o tamanho deste prefixo, o qual deve ser menor do que a máscara padrão da classe original do endereço. Deste modo, qualquer endereço pode ser representado pela forma prefixo/tamanho. Por exemplo, um endereço classe C – 192.32.1.0 – pode ser representado por um endereço CIDR – 192.32.1.0/24. Um endereço CIDR é denominado de agregado, pois contém vários endereços para uma mesma rota. Existe mais de um agregado para um endereço IP de acordo com o tamanho do prefixo adotado (192.32.0.0/16 ou 192.32.0.0/22, para o exemplo anterior). Quanto maior o tamanho do prefixo mais específico é o agregado.

O esquema CIDR provê um mecanismo de agregação das informações de rotas. Essa política de agregação de endereços diminui as tabelas de roteamento, mas alguns cuidados devem ser tomados para se evitar a introdução de rotas cíclicas [47]. Como uma tabela de rotas pode possuir mais de um agregado para um mesmo destino, porém passando por roteadores diferentes, a escolha do agregado é feita pela regra do maior prefixo, ou seja, a rota escolhida é aquela associada ao prefixo mais específico.

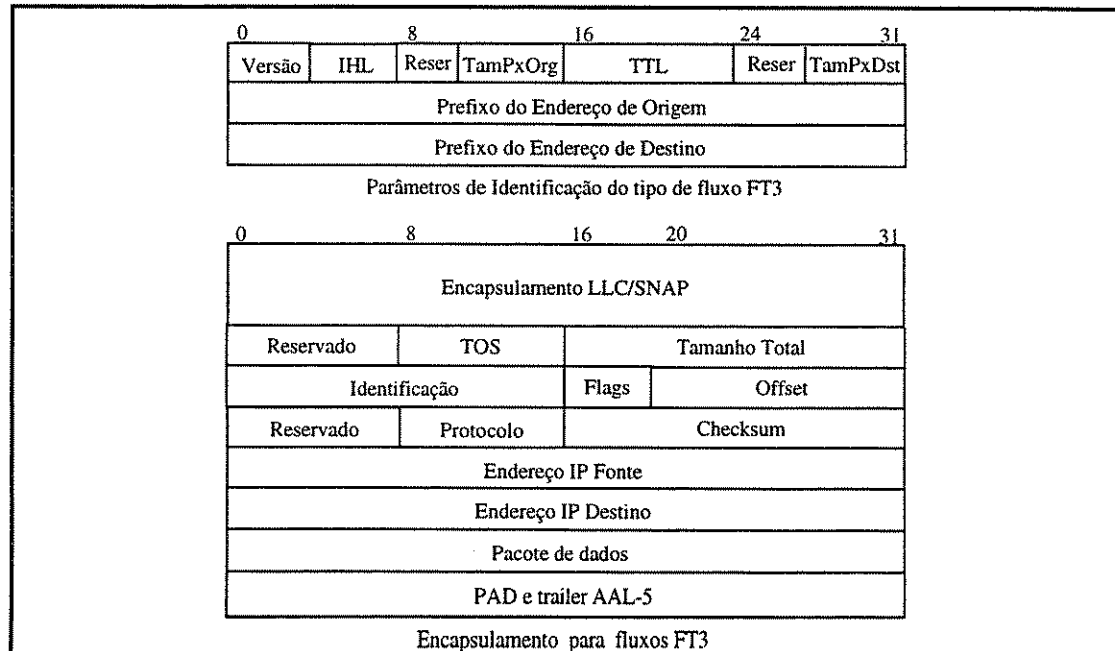


Figura 5.9: Definição do tipo de fluxo FT3

Outra diferença em relação ao encapsulamento dos outros tipos de fluxos é o envio em cada pacote dos endereços IP da origem e do destino, já que mais de um fluxo entre máquinas diferentes utiliza uma mesma conexão. Os parâmetros de identificação contém os prefixos de endereços CIDR de origem e destino e os tamanhos destes prefixos.

A necessidade do envio dos endereços existe pois as variáveis de estado só armazenam os prefixos de endereços, enquanto que os pacotes devem possuir endereços IP específicos. Desta forma, o encapsulamento habilita que opções de segurança (*firewalls*) possam ser estabelecidas. Uma máquina na fronteira de uma rede de rótulos pode testar os endereços dos pacotes recebidos com a máscara definida e então autenticar o pacote. Outra opção de segurança é a definição de um prefixo restritivo, ou seja, que só permita receber pacotes de um grupo IP específico.

Uma pequena modificação no objeto RSVP_LABEL deve ser realizada para suportar as variáveis de estado dos fluxos FT3 nos estilos FF e SE. Nesses estilos, as fontes precisam ser explicitamente indicadas nas mensagens RESV. Por isso, o prefixo dos endereços fonte é substituído pelos endereços utilizados nas mensagens RESV.

células desses fluxos são enviadas de forma intercalada pois não há qualquer controle de envio e em seu destino não é mais possível separá-las corretamente.

Para solução deste problema, a Comutação IS propõe a ordenação dos quadros AAL5 em filas nas portas de saída divididas por VPI/VCI de entrada. Com esse processo, todas as células de um quadro AAL5 proveniente de um canal virtual são armazenadas em sua respectiva fila até que a última célula seja recebida (a última célula de um quadro AAL5 é marcada com um bit especial no cabeçalho ATM). No momento em que todo quadro estiver agrupado, todas as células serão enviadas de forma contínua pela porta de saída. Não há possibilidade de intercalação de fluxos de canais virtuais distintos com esse esquema de enfileiramento. As filas podem ser implementadas em memória compartilhada de modo que as operações de adição e remoção de fluxos aos canais agregados sejam realizadas por simples movimentação de ponteiros sem gasto adicional de tempo ou sobrecarga do comutador.

O uso de comutadores com filas nas interfaces de saída acarreta atraso nos fluxos de células pois é necessário agrupar todo quadro AAL5. Esse processo aumenta o tamanho necessário da memória de armazenamento de células e transforma o padrão de tráfego aumentando as rajadas devido às transferências contínuas. As influências desses fatores se tornam mínimas se o tráfego tratado pelo comutador é muito alto [46]. Neste caso, devido à grande concorrência da interface de saída, o aumento do atraso e das rajadas causados pelo tempo de espera imposto às primeiras células do quadro até que a última esteja disponível se torna desprezível [46].

Dados mostram que, se a carga normal de operação de um comutador aumenta acima de 70%, não há aumento significativo do atraso nem das rajadas [46]. Essa carga é compatível com os comutadores atuais. O aumento do tamanho das filas é um fator mais restritivo porém com soluções factíveis.

5.6.5 Alteração no Tipo de Fluxo

Um Comutador IS com a característica de comutação adaptativa pode operar com várias funções na rede. O equipamento de ligação para redes periféricas ao núcleo da rede

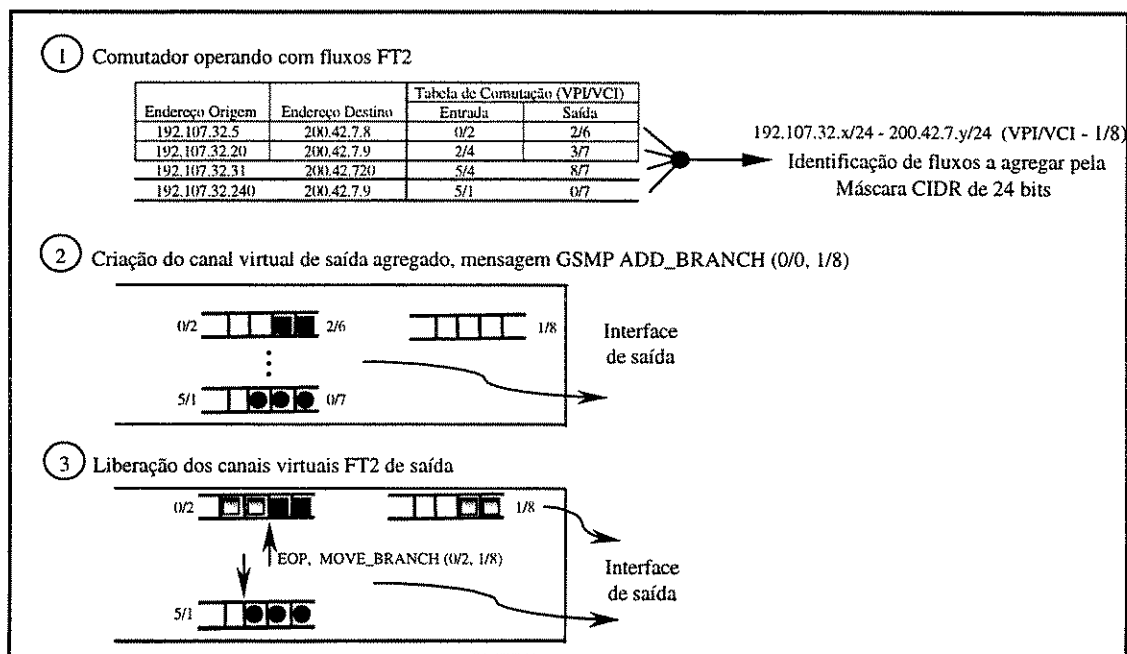


Figura 5.11: Mudança consistente de política de agregação

5.6.6 Análise da Agregação de Fluxos

Esta seção apresenta os resultados obtidos com a simulação de um comutador que utiliza as políticas de agregação definidas para a Comutação IS. Como discutido na proposta da arquitetura de Comutação IS, fluxos que necessitam de controle de QoS são associados a canais virtuais próprios para que possam receber escalonamento prioritário. Essa operação implica no aumento do número de canais virtuais usados no comutador, aproximando-se dos resultados obtidos com a política de fluxos FT1. A quantidade de canais virtuais disponíveis torna-se crítica com o aumento da demanda de tráfego, podendo esgotar o espaço disponível de rótulos na tecnologia ATM.

Como visto na seção anterior, a variação dos parâmetros X, Y e FTO resultam em um ponto de compromisso entre os resultados de canais virtuais gastos, a quantidade de eventos e a porcentagem de comutação. As políticas de agregação definidas pela arquitetura de comutação IS têm o objetivo de diminuir o número de canais virtuais necessários, aumentar a porcentagem de comutação e diminuir a frequência interna de gerência do comutador.

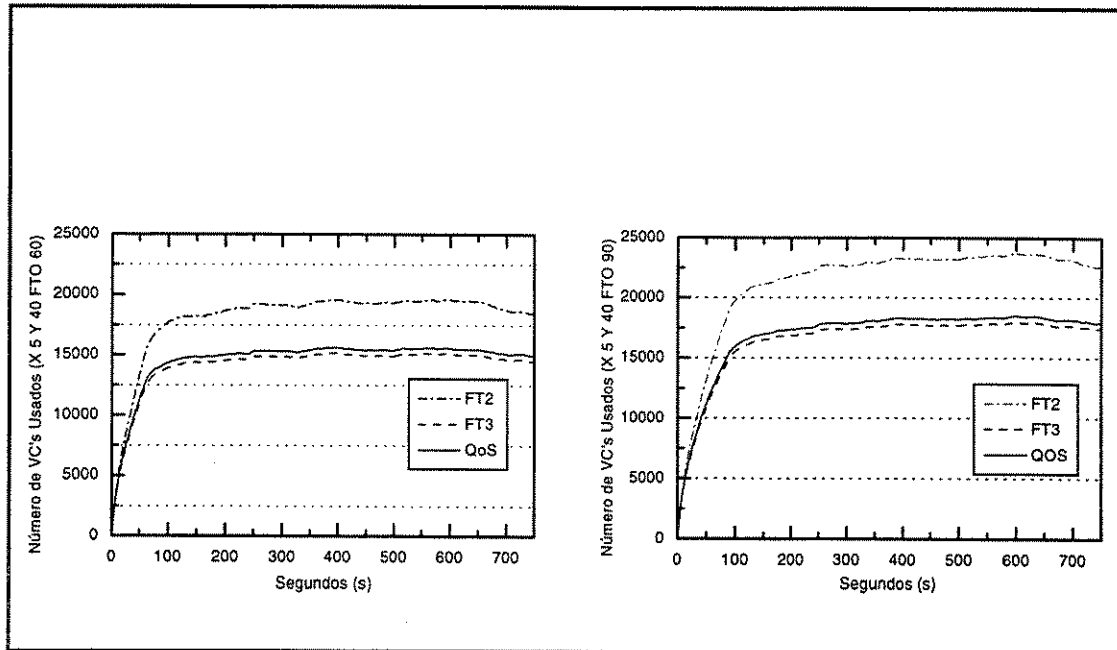


Figura 5.12: Comparação do uso dos tipos de fluxos FT2 e FT3

A Tabela 5.3 mostra os resultados obtidos para a porcentagem de pacotes comutados e o número médio de eventos gerados no comutador. A estratégia de agregação utilizada foi a de tipo de fluxo FT3. Esses números podem ser comparados com os valores obtidos nas seções anteriores com os mesmos parâmetros de operação do comutador. Deve-se notar a redução do número de canais virtuais e o aumento da porcentagem de pacotes comutados.

Crítérios de Classificação	Número de VC's Usados / s (média)	Pacotes Comutados (%)	Total de Eventos / s (média)
X 5 Y 30 FTO 60	16.621	–	–
X 5 Y 40 FTO 30	15.960	85,32	185,5
X 5 Y 40 FTO 90	19.200	87,35	137,4

Tabela 5.3: Comparação da variação dos parâmetros do Comutador IS (FixWest1)

A Figura 5.13 mostra o número total de eventos que ocorrem no comutador (ativações e desativações de conexões) para valor de FTO igual a 10 e 90 segundos. A política de agregação diminui a dinâmica de modificações nas tabelas internas do comutador pois as conexões permanecem por mais tempo abertas. Esta diminuição

pacotes. Deste modo, a média efetiva de utilização das conexões com política FT3 é maior pois as conexões FT4 permanecem mais tempo abertas e enviam um número não muito maior de pacotes comutados.

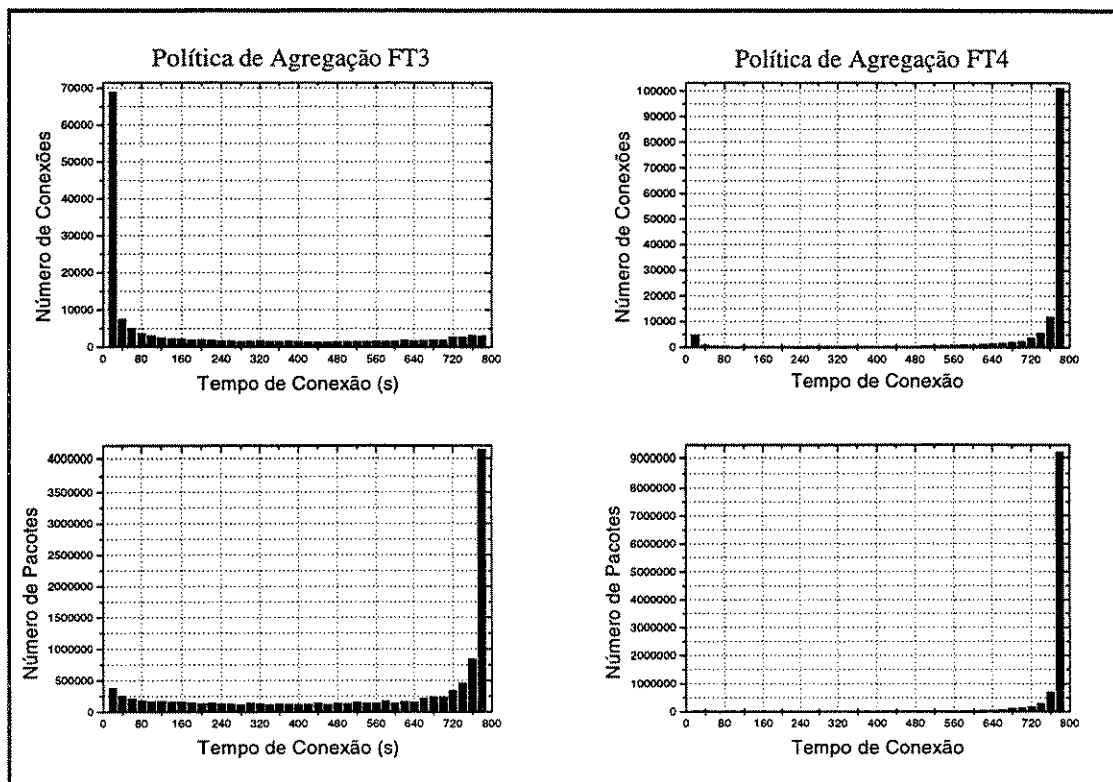


Figura 5.14: Utilização efetiva de conexões (distribuição de pacotes $\times$ tempo abertura)

O gráfico de agregação mostra que a maioria das conexões FT4 permanecem abertas a maior parte do tempo de simulação. Este comportamento se assemelha ao das conexões orientadas por topologia que ficariam sempre abertas. Porém, nem todas as rotas possuem tráfego. São aproximadamente 100.000 conexões abertas. Em uma rede pública isso representaria um gasto excessivo de recursos. As conexões FT3 que ficam abertas quase todo tempo de simulação é pelo menos uma ordem de grandeza menor. Como o arquivo de tráfego utilizado é o mesmo nos dois casos (mesma quantidade de pacotes), e como os resultados de simulação demonstram que os desempenhos de comutação são próximos, mostra-se que as conexões FT4 permanecem mais tempo ociosas pois quase a mesma quantidade de pacotes é comutada em canais FT3 que ficam em média ativas por um tempo menor.

Capítulo 6

Conclusão Geral

6.1 Considerações Finais

A tecnologia ATM está sendo amplamente utilizada para interligação de redes locais e de grandes distâncias. As redes de grandes distâncias foram o núcleo das redes de transportes de informações integradas. O desenvolvimento desta tecnologia foi rápido devido ao suporte dado pela indústria. Como fatores de sua disseminação estão a integração de voz, dados e vídeo em um único meio de transporte e a oferta de QoS na transmissão dessas informações.

Paralelamente, a família de protocolos TCP/IP é suportada sobre diversas arquiteturas e tornou-se um padrão de fato para protocolos de redes locais e de grandes distâncias. O modelo IIS foi padronizado para introduzir novos serviços IP e suportar tráfego sensível a requisitos de tempo. Seu uso está sendo consolidado com o crescimento das aplicações multimídia e sistemas de tempo real.

Nos últimos anos, tanto a indústria como a comunidade científica empregaram um grande esforço para desenvolver e melhorar mecanismos que permitissem uma melhor adaptação entre as duas arquiteturas. A arquitetura de Comutação por Rótulos é o mais recente esforço para incorporar à arquitetura ATM características inerentes aos protocolos IP, como flexibilidade e maior robustez a falhas.

A comutação por rótulos e o modelo IIS, padronizados pelo IETF, começam a ser empregados em conjunto com a tecnologia ATM. Porém, problemas intrínsecos da arquitetura ATM dificultam o oferecimento de serviços com QoS por esses modelos.

A arquitetura de Comutação IS foi proposta [62] com o objetivo de otimizar a integração dos mecanismos do modelo IIS e o protocolo RSVP com comutadores baseados na tecnologia ATM e que utilizam a comutação por rótulos para envio integrado das informações. Desta forma, o envio de dados pela camada de enlace é otimizado no

do comutador. Porém, essas alterações causam impacto negativo na carga de gerência de recursos do comutador ou na porcentagem de comutação.

Baseado nos resultados obtidos na simulação, a Comutação IS foi proposta com um mecanismo que se adapta ao padrão de tráfego sustentado e aos vários tipos de fluxos existentes na rede. Isso possibilita que o Comutador IS suporte as mudanças de padrões de tráfego comuns na Internet e o surgimento de tráfego sensível a tempo.

O desempenho da nova arquitetura foi investigado em simulação que validou os resultados obtidos quando comparados com os obtidos na literatura. O problema de escassez de rótulos pode ser reduzido pela variação dos parâmetros do comutador, mas essas mudanças exercem impacto negativo na porcentagem de pacotes comutados e na frequência de atualização das tabelas de comutação.

Os resultados de simulação também mostraram que o desempenho obtido com as políticas de agregação foram positivos pois tanto a gerência de recursos como a porcentagem de comutação se mostraram em valores satisfatórios. A agregação de canais virtuais é fundamental para conexões compartilhadas pois reduz a escassez de rótulos e melhora a comutação de pacotes e frequência de atualização.

O desempenho de comutação também aumenta com a utilização de políticas de maior agregação. Outro fator positivo é a redução de eventos no comutador devido a diminuição de criação de canais (conexões mais estáveis). Desta forma, o mecanismo de agregação torna-se vantajoso para o núcleo da rede onde diversos fluxos podem se beneficiar de um canal agregado.

6.2 Cenário de Utilização

A Figura 6.1 propõe uma rede multi-serviços baseada na comutação de fluxos agregados por topologia. Essa rede é formada por um conjunto de comutadores convencionais associados a um roteador clássico que controla a comutação na camada de enlace. O roteador pode enviar dados de maneira convencional por uma conexão preestabelecida nos comutadores. Porém, o controle no roteador pode definir caminhos comutados para enviar fluxos selecionados por um critério que se adapta às mudanças de

6.3 Contribuições e Trabalhos Futuros

Este trabalho teve como objetivo inicial estudar e investigar os problemas decorrentes da integração da tecnologia ATM com o modelo *IIS*. Este estudo possibilitou o desenvolvimento de uma nova arquitetura de comutação que integra fortemente os mecanismos do modelo *IIS* com a tecnologia ATM.

Esta nova arquitetura foi denominada de Comutação IS (*IntServ*). A idéia é que a comutação nesta arquitetura seja flexível quanto ao estabelecimento de conexões e que se adapte às mudanças nos padrões de tráfego. Foi proposto neste trabalho um mecanismo de agregação de fluxos que permite poupar canais virtuais, suportar sessões compartilhadas RSVP e tráfego com QoS.

Como estratégia de controle do uso de canais virtuais e para possibilitar o oferecimento de conexões com QoS, a Comutação IS foi proposta com dois tipos de fluxos adicionais que permitem a agregação do tráfego de diferentes canais virtuais: FT3 (fluxo entre um par de sub-redes) e FT4 (fluxo para uma sub-rede ou com destino a um roteador de saída do domínio).

Com a arquitetura de Comutação IS, mesmo com o uso de tipos de fluxos agregados que se aproximam da operação por topologia, é possível utilizar as características de orientação a tráfego, restringindo problemas como a manutenção de conexões ociosas e a grande quantidade de conexões abertas em redes com muitos membros (explosão de rotas).

Foi proposto um mecanismo que permite ao Comutador IS operar em modo que se adapta ao padrão de tráfego utilizado. A arquitetura proposta foi validada através de simulação. Esta simulação permitiu realizar comparações com resultados obtidos pela literatura com outras arquiteturas. Os resultados obtidos na simulação indicam que a agregação trouxe bons resultados.

O mecanismo de agregação liberou conexões para serem oferecidas para fluxos com QoS e suporte aos estilos compartilhados RSVP (conexões multiponto–multiponto).

Referências Bibliográficas

- [1] D. Clark, S. Shenker, L. Zhang, “*Supporting Real-time Applications in an Intergrated Services Packet Network: Architecture and Mechanisms*”, Proc. ACM SigCOMM'92, Baltimore, Agosto 1992.
- [2] D. Clark, S. Shenker, R. Braden, “*Integrated Services in Internet Architecture: an Overview*”, IETF RFC 1633, Julho 1994.
- [3] S. Blake, D. Black, et al, “*An Architecture for Differentiated Services*”, IETF RFC 2475, Dezembro 1998.
- [4] S. Shenker, L. Breslau, “*Two Issues in Reservation Establishment*”, Proc. ACM SigCOMM'95, Cambridge, Agosto 1995.
- [5] S. Shenker, “*Fundamental Design Issues for the Future Internet*”, IEEE JSA in Communications, págs. 1176-1188, Setembro 1995.
- [6] L. Zhang, S. Deering, et al., “*RSVP: a new Resource ReSerVation Protocol*”, IEEE Network Magazine, págs. 8-18, Setembro 1993.
- [7] R. Braden, L. Zhang, et al., “*Resource ReSerVation Protocol (RSVP) - version 1 Funcional Especification*”, IETF RFC 2205, Setembro 1997.
- [8] P. White, “*RSVP and Integrated Services in the Internet: A Tutorial*”, IEEE Communications Magazine, págs. 100-106, Maio 1997.
- [9] J. Wroclawski, “*The use of RSVP with IETF Integrated Services*”, IETF RFC 2210, Setembro 1997.
- [10] S. Shenker, J. Wroclawski, “*Network Element Service Specification Template*”, IETF RFC 2216, Setembro 1997.

- [23] M. Garrett, M. Borden, "*Interoperation of Controlled-Load and Guaranteed Service with ATM*", IETF RFC 2381, Agosto 1998.
- [24] R. Callon, P. Doolan, et al., "*A Framework for Multiprotocol Label Switching*", IETF DRAFT MPLS WG, Setembro 1999.
- [25] E. Rosen, A. Viswanathan, R. Callon, "*Multiprotocol Label Switching Architecture*", IETF DRAFT MPLS WG, Agosto 1999.
- [26] A. Viswanathan, N. Feldman, et al., "*Evolution of Multiprotocol Label Switching*", IEEE Communications Magazine, págs. 165-173, Maio 1998.
- [27] S. Keshav, R. Sharma, "*Issues and Trends in Router Design*", IEEE Communications Magazine, págs. 144-151, Maio 1998.
- [28] P. Newman, G. Minshall, et al., "*IP Switching and Gigabit Routers*", IEEE Communications Magazine, págs. 64-69, Janeiro 1997.
- [29] _____, "*Flow Labelled IP: A Connectionless Approach to ATM*", Proc. IEEE Infocom, San Francisco, Março 1996.
- [30] _____, "*Ipsilon Flow Management Protocol Specification for IPv4*", IETF RFC 1953, Maio 1996.
- [31] _____, "*Ipsilon General Switch Management Protocol Specification – version 1.1*", IETF RFC 1987, Maio 1996.
- [32] _____, "*Ipsilon General Switch Management Protocol Specification – version 2.0*", IETF RFC 2397, Março 1998.
- [33] _____, "*Transmission of Flow Labelled IPv4 on ATM Data Link*", IETF RFC 1954, Maio 1996.
- [34] T. Worster, "*A QoS Model for GSMP*", IETF DRAFT GSMP WG, Fevereiro 1999.

- [46] A. Schmid, I. Iliadis, P. Droz, “*Impact of VC Merging on Buffer Requirements in ATM Networks*”, Conference on High Performance Networking (HPN’98 – Viena), Setembro 1998.
- [47] V. Fuller, T. Li, et al., “*Classless Inter-Domain Routing (CIDR): an Address Assignment and Aggregation Strategy*”, IETF RFC 1519, Setembro 1993.
- [48] P. Potter, M. Zukerman, “*Analysis of a Discrete Multipriority Queueing System Involving a Central Shared Processor Serving Many Local Queues*”, IEEE JSA in Communications, págs. 192-202, fevereiro 1991.
- [49] M. Garret, “*A Service Architecture for ATM: From Applications to Scheduling*”, IEEE Networks, págs. 6-15, Junho 1996.
- [50] S. Keshav, “*An Engineering Approach to Computer Networking: ATM Networks, the Internet and the Telephone Network*”, Addison-Wesley, cap. 9, págs. 209-263, Edição 1998.
- [51] W. Fischer, E. Wallmeier, et al., “*Data Communications Using ATM: Architectures, Protocols, and Resource Management*”, IEEE Communications Magazine, págs. 24-33, Agosto 1994.
- [52] N. Finn, T. Mason, “*ATM Lan Emulation*”, IEEE Communications Magazine, págs. 96-100, Junho 1996.
- [53] M. Laubach, J. Halpern, “*Classical IP and ARP over ATM*”, Internet RFC 2225, Abril 1998.
- [54] D. McDysan, D. Spohn, “*ATM: Theory and Application*”, McGraw-Hill, New York, Edição 1996.
- [55] D. Comer, “*Internetworking with TCP/IP, vol. I*”, McGraw-Hill, New York, Edição 1995.

Apêndice A

Interconexão de Redes ATM

A.1 Introdução

As redes digitais de serviços integrados de faixa larga – RDSI/FL (*Broadband - Integrated Services Digital Networks*, B/ISDN) – surgiram com o propósito de integrar a transmissão de voz e vídeo com o envio de dados. A tecnologia ATM foi escolhida como base para implementação das RDSI/FL por comitês internacionais como: ATM Fórum, ITU-T (*Internacional Telecommunications Union – Telecommunications sector*), ANSI (*American National Standards Institute*). As redes ATM são um conjunto de mecanismos e protocolos para envio de informações e para troca de sinalização propostos e padronizados por esses comitês [60].

A tecnologia ATM é orientada à conexão [51], isto é, os dados só são transferidos depois que canais virtuais de comutação sejam explicitamente estabelecidos pela configuração de circuitos entre a fonte e o receptor, semelhante às redes públicas de telefonia. São utilizados rótulos para a comutação dos dados por esses canais virtuais. As redes ATM tem a característica de não suportar difusão de mensagens (*NonBroadcast Multi-Access Networks*) na sua forma nativa, e necessita para isso recorrer a mecanismos complexos para envio em grupos (conexões ponto–multiponto, servidores de difusão).

Existem dois principais tipos de soluções para a adaptação dos protocolos da camada de rede à tecnologia ATM. No primeiro modo, denominado de sobreposição de camadas (ou simplesmente, sobreposto), todo o esquema de endereçamento, roteamento e outras funcionalidades construídas pelas camadas ATM são sobrepostas pela funcionalidade IP de forma independente. No segundo modo, conhecido como nativo, mecanismos de resolução de endereços são utilizados para resolver endereços de rede diretamente em endereços ATM, e desta forma, os pacotes da camada de rede são carregados diretamente em conexões ATM [16].

um único servidor e todo componente que quiser pertencer a rede deve ser controlado por ele. Especifica aos clientes como se conectar a rede lógica local.

Servidor de difusão (BUS - *Broadcast and Unknown Server*): realiza o serviço de enviar uma cópia de um pacote a todas as máquinas de uma ELAN, ou seja, implementa o serviço de *multicast*. É utilizado também para o envio de pacotes a uma máquina com endereço ainda desconhecido.

Servidor de configuração (LECS - *LAN Emulation Configuration Server*): sua função é atribuir um servidor LANE a um cliente que requisite sua entrada em uma ELAN. Existe um único LECS servindo várias redes emuladas dentro de um único domínio administrativo.

Os vários componentes de uma ELAN se comunicam entre si através de uma série de conexões ATM. Essas conexões são separadas para a transmissão de dados e de tráfego de controle. As conexões de controle são:

- VCC (*Virtual Channel Connection*) direto de configuração: conexão bidirecional ponto a ponto de todo LEC até um LECS para obtenção de dados de inicialização.
- VCC direto de controle: conexão bidirecional ponto a ponto estabelecido de todo LEC até um LES para troca de mensagens de controle.
- VCC distribuído de controle: conexão ponto-multiponto unidirecional ligando um LES a todos os clientes LEC de uma rede emulada.

Os canais de conexão destinados a transmissão de dados definidos pelo modelo LANE são os seguintes:

- VCC direto de dados: conexão bidirecional ponto a ponto estabelecida entre dois clientes LEC para troca de dados, geralmente é utilizado uma única conexão entre dois endereços MAC com o objetivo de diminuir o atraso e a utilização de recursos.
- VCC de envio de difusão: conexão bidirecional ponto a ponto ligando o LEC ao BUS para envio de dados com o objetivo de transmissão *multicast*.

A.3 IP Clássico sobre ATM

A família de protocolos TCP/IP foi inicialmente projetada para interligar redes não orientadas à conexão e com suporte a difusão de dados. Uma rede ATM não possui tais características portanto houve a necessidade da definição de métodos específicos [53] que possibilitassem a adaptação da arquitetura IP a essa nova tecnologia. Esses métodos introduzem alterações no encapsulamento de dados e no protocolo de resolução de endereços que são utilizados nas redes tradicionais.

A.3.1 Encapsulamento de Pacotes

A tecnologia ATM oferece a camada de adaptação AAL-5 (*ATM Adaptation Layer 5* - Figura A.2 [54]) para transporte de dados, permitindo o envio de pacotes de tamanho variável, entre 1 a 65.535 octetos, em células de tamanho fixo através de segmentação e reconstrução dos pacotes de dados. O IP clássico utiliza a camada AAL-5 para transmissão dos seus pacotes restringindo o tamanho máximo com o estabelecimento do valor da MTU (*Maximum Transfer Unit*) em 9180 octetos.

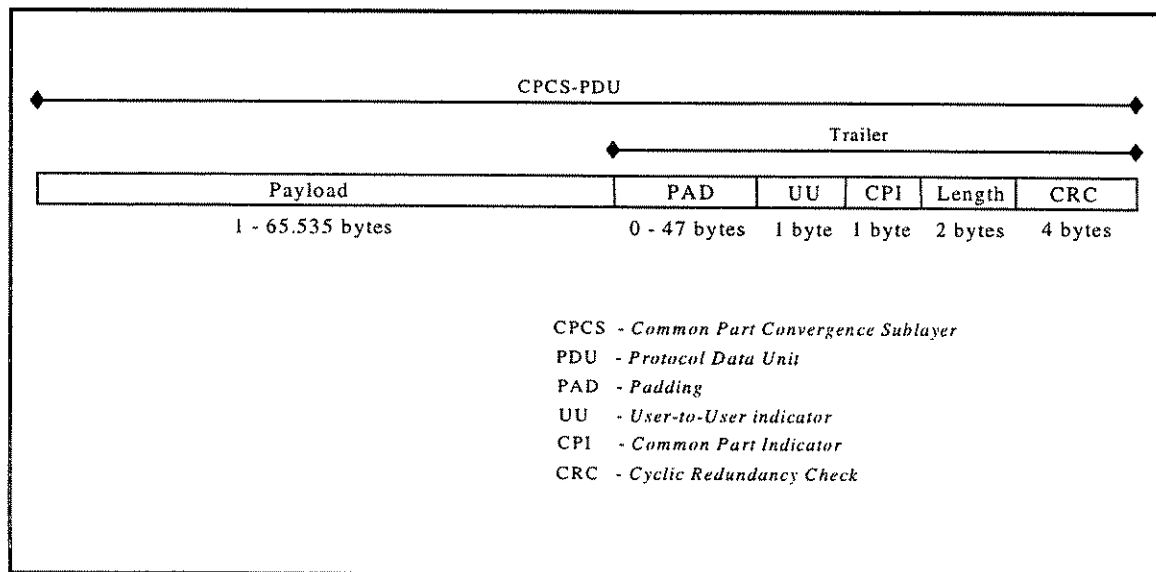


Figura A.2: Formato do quadro AAL5

(*ATM Address Resolution Protocol*) armazena as associações de endereços IP e ATM de todas as máquinas pertencentes a uma LIS, e sempre que for necessário o estabelecimento de uma conexão a um endereço ATM desconhecido, uma requisição de resolução deve ser enviada a este servidor que responderá com algum endereço armazenado correspondente a outra máquina da LIS. Depois de obtido o endereço ATM do próximo nó de comunicação, um circuito virtual pode ser estabelecido através do uso dos procedimentos de sinalização UNI (*User Network Interface*) determinados para o IP clássico [21].

O gerenciamento de conexões se resume ao controle do tempo em que um circuito virtual é mantido aberto. Como o procedimento de estabelecimento de um circuito virtual é demorado, se um fluxo de pacotes IP atravessar uma conexão de um modo contínuo deve se optar pela manutenção da abertura do canal, caso contrário o canal deve ser fechado pois ficará ocioso por grandes períodos de tempo, representando um custo adicional.

A.4 Modelo baseado em nuvens ATM

O modelo desenvolvido em [56,57] tem a visão de uma rede ATM como uma grande nuvem em que todas as máquinas pertencentes a ela podem estabelecer conexões diretamente, independentemente da maneira como estão divididas em subgrupos lógicos (LIS). Esta nuvem tem a característica de permitir múltiplos acessos de máquinas a rede, porém sem a existência de um suporte próprio a difusão de mensagens. Deste modo, foram desenvolvidos, em paralelo, protocolos independentes para suporte a difusão de mensagens e a resolução de endereços.

A.4.1 Redes com Acesso Múltiplo sem Difusão

O modelo é baseado no protocolo NHRP [57] (*Next Hop Resolution Protocol*) e foi construído tendo como base o modelo IP clássico, mas com o objetivo de acabar com as limitações sofridas por este. Neste modelo uma máquina pode estabelecer uma conexão direta com qualquer outra máquina fisicamente ligada a nuvem ATM, ao

servidores também armazenem aquele endereço. De posse do endereço requisitado, o primeiro NHS envia ao cliente uma mensagem de resposta NHRP (passo 3 da figura acima). Uma conexão de dados é então estabelecida (em nível ATM com a sinalização Q.2931) ligando diretamente as máquinas envolvidas na transmissão de dados, sendo possível a utilização de QoS ATM.

Devido ao tempo de espera que pode ser gerado com a resolução de endereços pelos servidores NHS somado ao tempo de estabelecimento da conexão na rede ATM, o protocolo sugere que pacotes urgentes possam ser enviados por um caminho padrão até que a conexão não esteja disponível. Porém isso pode gerar uma degradação no desempenho do sistema devido a chegada de pacotes fora de ordem.

A.4.2 Servidores de Difusão

Com o objetivo de possibilitar a comunicação por difusão de mensagens e cobrir a falta deste tipo de transmissão no modelo baseado em LIS, básica para o funcionamento da arquitetura IP, foi proposto o protocolo MARS (*Multicast Address Resolution Server*) [58]. Esse protocolo também pode ser utilizado no modelo de redes sem difusão não apresentando dificuldades de adaptação.

Sua operação é dividida em dois casos de acordo com o modo de transmissão para o grupo *multicast*. No primeiro modo, denominado servidor, o servidor MARS envia ao nó uma lista de servidores de envio de difusão que estão conectados ao grupo *multicast*, neste caso o nó transmite os dados para esses servidores que as repassam a todo grupo. No segundo modo, denominado malha, o servidor MARS envia o endereço de todas as máquinas do grupo *multicast* diretamente ao nó, este então estabelece conexões ATM diretamente com todas as máquinas do grupo, cabendo ao servidor avisar sobre eventuais mudanças no grupo para o gerenciamento de conexões realizado pelo nó fonte.

Cientes MPOA: são máquinas ligadas diretamente a rede ATM e que conhecem todo o procedimento MPOA, negociando suas requisições diretamente com os servidores.

Elementos de borda e Cliente MPOA são agrupados em conjuntos denominados de IASG (*Inteworking Address SubGroup*) que são equivalentes a uma LIS IP ou a um grupo ELAN. Coordenadores de IASG distribuem todos os clientes que requisitam entrada na rede MPOA em algum grupo IASG, executando a mesma função que um servidor de configuração LANE.

Como no modelo LANE, a arquitetura MPOA define servidores de difusão e envio padrão, cujas funções são de determinação de um endereço desconhecido por difusão e o envio de dados enquanto uma associação de endereços ainda está sendo resolvida.

A arquitetura MPOA se utiliza do mecanismo ATMARP definido pelo IP clássico para realizar a associação entre endereços de rede e endereços físicos ATM. Porém seus servidores de rota se baseiam no protocolo NHRP para realizar de forma distribuída a busca pela rede MPOA de um endereço requisitado.

Em adição a essas características, a solução MPOA suporta o conceito de roteamento virtual. Um roteador virtual é definido como um conjunto de servidores dispersos geograficamente que executam a função de roteamento dividida em dois módulos: (1) de controle, que mantém informações de roteamento e executa os protocolos para computação das rotas; e (2) interface, que armazena as rotas a fim de responder as requisições dos clientes.

Apêndice B

Estrutura de Simulação

B.1 Introdução

Neste trabalho, a simulação foi utilizada com o objetivo de validar a proposta da arquitetura de Comutação IS e comparar os resultados obtidos com os disponíveis em trabalhos anteriores de outras arquiteturas .

Foi utilizada a ferramenta *issim*¹ para simulação de eventos discretos em um comutador com orientação por fluxos. As características de agregação e reconhecimento de fluxos prioritários foram introduzidas por alterações na estrutura do simulador.

A seção B.2 explica a estrutura original do simulador implementado na literatura. A seção B.3 discute os arquivos de tráfego utilizados na simulação: onde foram obtidos, sua estrutura e topologia da rede. A seção B.4 apresenta as modificações implementadas no simulador original de modo a acrescentar características da Comutação IS.

B.2 Arquitetura de simulação

O simulador *issim* [44] foi desenvolvido com o objetivo de obter resultados práticos do funcionamento da estrutura de controle de um comutador que se utiliza da arquitetura *IP Switch*. O código fonte foi escrito em linguagem C (padrão ANSI) para plataforma UNIX.

A estrutura de funcionamento é apresentada na Figura B.1, onde são vistos os possíveis caminhos em que as células contendo os pacotes percorrem dentro do comutador e o tempo gasto até seu envio pela porta de saída. As células ATM que carregam pacotes de fluxos já classificados são enviadas diretamente para a porta de saída

¹ O *IP Switch Simulator* foi obtido em <ftp://klamath.stanford.edu/pub/outgoing/issim.tar.gz>

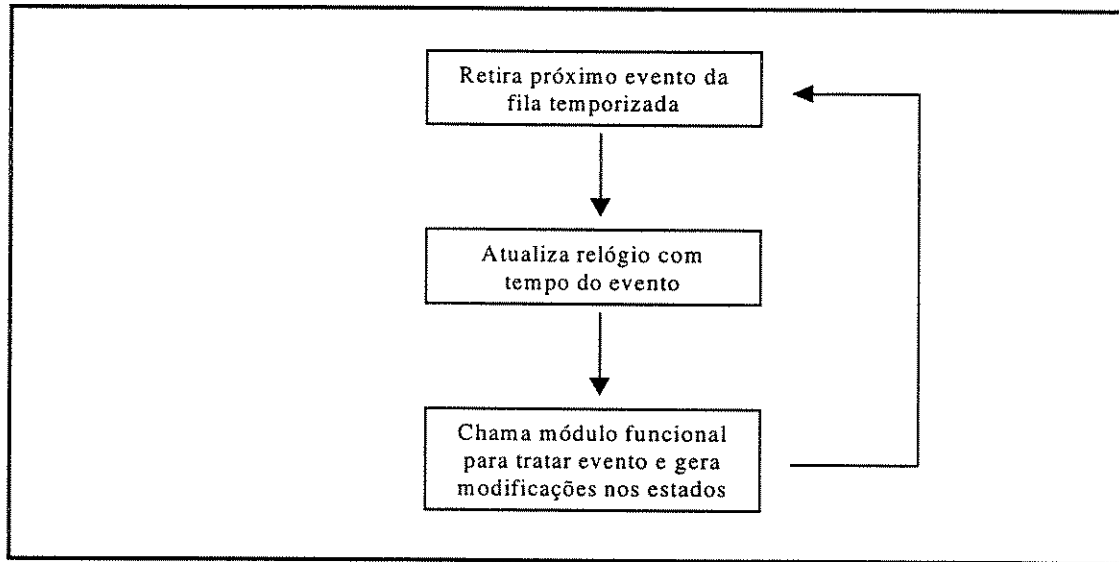


Figura B.2: Ciclo de simulação de eventos

A Figura B.3 apresenta o diagrama funcional de simulação, onde o evento de chegada de um pacote inicia a operação do simulador. O módulo funcional de chegada submete o pacote à decisão de comutação consultando as bases de dados do comutador, em particular a tabela de fluxos comutados que contém os dados de comutação VPI/VCI da camada de enlace (e que será apresentada nas próximas figuras).

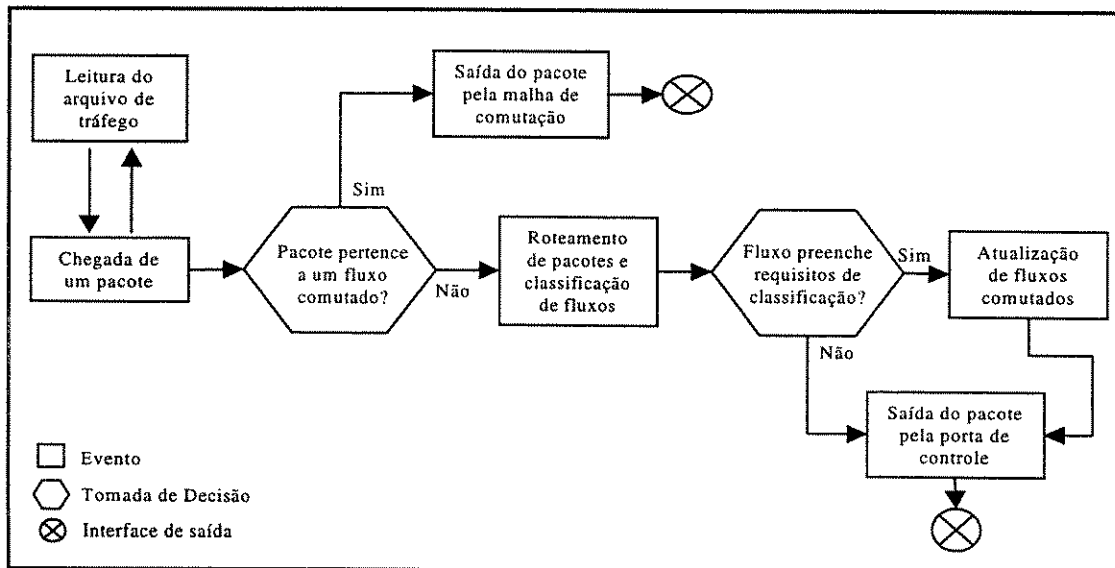


Figura B.3: Diagrama Funcional do simulador *issim*

As bases de dados são organizadas na forma de tabelas *hash*, como mostrado na figura. Existem duas tabelas: (1) a tabela de fluxos comutados que armazena as variáveis de estado e temporização de fluxos já classificados, e (2) a tabela de fluxos candidatos que armazena as variáveis de estado, temporização e a quantidade de pacotes dos fluxos que estão sendo classificados.

O simulador pode utilizar três métodos para classificação de fluxos: (1) portas (http, ftp, etc.); (2) tipo de protocolo (UDP/TCP); ou (3) intensidade do fluxo (através de taxas X/Y , X – número de pacotes recebidos em Y – período de tempo). Apenas o classificador X/Y é considerado para análise neste trabalho, pois seu desempenho é sempre melhor do que os dos outros métodos quando se analisa o número de canais virtuais alocados [44].

No método de classificação X/Y , as bases de dados são utilizadas para consultas pelos módulos funcionais dos eventos para decisão de comutação ou não dos pacote, sendo atualizadas quando um fluxo atinge os critérios determinados.

B.3 Arquivos de Tráfego

Os arquivos de tráfego (*tracefiles*) utilizados na simulação são medidas reais de interfaces de redes em operação, de onde os cabeçalhos de pacotes IP foram capturados pelo utilitário *snoop* do sistema *Sun Solaris* [61]. Os dois primeiros arquivos (*DecPkt*²) são medidas de uma rede local corporativa, enquanto o último (*FixWest*³) foi coletado de um roteador central (*backbone*) de uma rede de pesquisas (*vBNS* – *Virtual Backbone Network Services*) na costa oeste dos Estados Unidos. A Figura B.5 mostra detalhes sobre a topologia da rede onde o tráfego foi capturado.

² Obtido no Internet Traffic Archive (<http://ita.ee.lbl.gov/html/contrib/dec-pkt.html>) em 2/5/98

³ Obtido no National Lab for Applied Network Research (<ftp://oceana.nlanr.net/Traces/FR+>) em 2/5/98

Os arquivos de tráfego obtidos na Internet trazem os endereços IP codificados por questões de segurança. Essa codificação mantém a integridade das informações de domínios pois não elimina a estrutura de topologia dos endereços IP. Os endereços após a codificação mantém a relação de domínio IP ao qual pertenciam, independente da classe específica de cada endereço. Essa manutenção da integridade de classe do endereço IP permite a análise coerente dos dados através de agregação por domínios. A Tabela B.1 mostra o programa na linguagem *perl* em que o arquivo *FixWest* foi codificado.

```
Codifica os endereços IP em 16 bits para prefixo de rede e 16 bits para
endereços de máquinas. A codificação usa as classes A, B, C e D/E do IP.

$infile=$ARGV[0]; open(infile,$infile) || die;

while (read(infile,$record,44)) {

    $src=vec($record,5,32);
    $sid=($src>>28)&0xf;

    $srcn=$src & 0xf0000000;$srch=$src & 0x0ffffff;
    if ($sid < 0xe) {$srcn=$src & 0xfffff00;$srch=$src & 0x000000ff;}
    if ($sid < 0xc) {$srcn=$src & 0xffff000;$srch=$src & 0x0000ffff;}
    if ($sid < 0x8) {$srcn=$src & 0xff00000;$srch=$src & 0x00ffffff;}
    if (!$id{$srcn,-1}) {$id{$srcn,-1} = ++$netcount;}
    if (!$id{$srcn,$srch}) {$id{$srcn,$srch} = ++$hstcount($srcn);}
    $src=($id{$srcn,-1}<<16) + $id{$srcn,$srch};
    vec($record,5,32)=$src;

    $dst=vec($record,6,32);
    $did=($dst>>28)&0xf;

    $dstn=$dst & 0xf0000000;$dsth=$dst & 0x0ffffff;
    if ($did < 0xe) {$dstn=$dst & 0xfffff00;$dsth=$dst & 0x000000ff;}
    if ($did < 0xc) {$dstn=$dst & 0xffff000;$dsth=$dst & 0x0000ffff;}
    if ($did < 0x8) {$dstn=$dst & 0xff00000;$dsth=$dst & 0x00ffffff;}
    if (!$id{$dstn,-1}) {$id{$dstn,-1} = ++$netcount;}
    if (!$id{$dstn,$dsth}) {$id{$dstn,$dsth} = ++$hstcount($dstn);}
    $dst=($id{$dstn,-1}<<16) + $id{$dstn,$dsth};
    vec($record,6,32)=$dst;

    $iprot=vec($record,17,8)
    if($iprot == 6){
        syswrite(stdout,$record,44);
    } elsif($iprot==17) {
        syswrite(stdout,$record,36)
        syswrite(stdout,$record,8);
    } else {
        syswrite(stdout,$record,28)
        syswrite(stdout,$record,16);
    }
}
```

Tabela B.1: Listagem do arquivo de codificação em *perl*

Os registros dos arquivos de tráfego representam a chegada de um pacote e são utilizados como dados de entrada para a simulação. Os registros são estruturados da seguinte forma: (1) tempo de chegada do pacote (*timestamp*), (2) o cabeçalho IP

características *Tag Switch* e *IP Switch*. Com isso, foi possível medir o tempo médio em que conexões orientadas por topologia e as orientadas por tráfego permanecessem abertas. Desta forma é possível comparar os requisitos dos dois tipos de arquiteturas.

A estrutura principal de simulação não foi alterada, sendo mantido todo mecanismo de implementação tempo real (como a fila cronológica de eventos e o mecanismo de atualização do relógio). As modificações no ambiente de simulação foram implementadas na estrutura dos módulos funcionais, principalmente com a adição do controle de fluxos baseado no tipo FT3/FT4 em substituição ao controle FT1/FT2.

As alterações foram realizadas nas bases de dados de comutação e nas rotinas para decisão de reconhecimento de fluxo. Nas bases de comutação, os registros de identificação de fluxo foram alterados para que pudessem armazenar os parâmetros de identificação de fluxos FT3/FT4. As rotinas de decisão foram mudadas para que possam ser aplicadas máscaras de endereços ao cabeçalho dos pacotes IP, para que desta forma, os pacotes sejam agregados por topologia.

Apesar de fornecer a funcionalidade básica, o simulador utilizado como base é limitado. Não permite a utilização de um comutador com mais de uma porta de entrada e outra de saída. Além disso, a simulação é de apenas um comutador isolado, não havendo suporte para implementação de uma rede de comutadores. Essas características do simulador *issim* dificultaram a obtenção de resultados da interação entre comutadores e do processo de estabelecimento de rotas RSVP mais complexas.

Essas características não foram adicionadas ao simulador modificado por este trabalho. Ao invés disso, foram assumidas algumas situações e cenários para obtenção desses resultados. Como exemplo, foi assumido que o simulador se comportasse como uma máquina central de um *backbone* de uma rede metropolitana de acordo com a topologia de onde foi coletado o arquivo de tráfego. Como outro exemplo, foi a adoção de portas de transporte de áudio e vídeo existentes no arquivo de tráfego como fluxos prioritários RSVP. Não há tráfego RSVP já existente nas redes já em operação.