

Avaliação de Desempenho de Serviço Multicast em Redes IP sobre ATM

¹Ricardo Yoshio Horita, ²Luiz E.C.Bergamo, ³Luis Carlos Trevelin

Universidade Federal de São Carlos
Departamento de Ciência da Computação
Rodovia Washington Luiz, km 235 São Carlos, Brasil
13.565 – POBOX 676
Fone: +55 016 260 8233
¹horita@salesianolins.br
²bergamo@salesianolins.br
³trevelin@dc.ufscar.br

Resumo—

ATM (*Asynchronous Transfer Mode*) é uma tecnologia que permite a transferência de dados multimídia em uma única conexão física, com elevadas taxas de transmissão (acima das tecnologias de redes existentes), possibilitando, ainda, reservar Qualidade de Serviço necessária para a transmissão de áudio e vídeo. O modelo IP Clássico sobre ATM busca adaptar o protocolo IP utilizado nas redes atuais à tecnologia ATM. A transmissão IP *multicast* sobre ATM permite otimizar a comunicação de grupos de usuários utilizando redes IP sobre ATM.

A transmissão de dados IP *multicast* sobre ATM utilizando UNI 3.0/3.1 pode ser feita através de dois sistemas: VC *Meshes* ou com servidor *multicast*.

Este trabalho tem por objetivo modelar e avaliar o desempenho das redes IP *multicast* sobre ATM utilizando UNI 3.0/3.1, buscando fazer uma análise comparativa entre a utilização de VC *Meshes* e com Servidor *Multicast* (MCS)

Palavras-chave— IP, *Multicast*, ATM, Servidor de Resolução de Endereços *Multicast* (MARS), VC *Meshes*, servidor *multicast* (MCS), Rede de Filas.

Abstract—

ATM (*Asynchronous Transfer Mode*) technology provides transfer multimedia data over a unique physical connection, with high transmission rates (above the existent network technology), allowing Quality of Service reservation, necessary to audio and video transmission. The Classical IP model over ATM fits IP protocol used in current networks to ATM technology. The multicast transmission IP over ATM improves the users's group communication by using IP multicast networks over ATM.

The Multicast IP data over ATM utilizing UNI 3.0/3.1 may be done through two systems: VC *Meshes* or Multicast Servers.

The aim of this study is to model and evaluate the performance of multicast IP over ATM networks employing UNI 3.0/3.1 and also make a comparative performance analysis between VC *Meshes* and the Multicast Servers.

Keywords—IP, Multicast, Asynchronous Transfer Mode Multicast (ATM), Multicast Address Resolution Server (MARS), Virtual Connections (VC) meshes, Multicast Server (MCS), Queueing Network.

I. INTRODUÇÃO

A tecnologia de rede ATM (*Asynchronous Transfer Mode*) permite a transmissão de dados multimídia em uma única conexão com reserva de Qualidade de Serviço [CAR 98][CAB 97] à taxa de transmissão muito mais elevada que as oferecidas pelas tecnologias de redes tradicionais existentes (IP, *frame relay*, etc).

Com o aparecimento de aplicações que demandam alto desempenho de comunicação, tais como videoconferência, educação à distância, trabalhos cooperativos, que necessitam de comunicação entre membros de um grupo via rede de computadores, surgiu a necessidade de se criar uma tecnologia que permitisse otimizar esta espécie de transmissão de dados. Denominada transmissão *multicast*, segundo Armstrong [AMS 92], esta comunicação de grupos permite desempenhar colaboração em tempo real através de aplicações colaborativas como projeto de um novo produto com utilização de *whiteboard* [KOS 98], sem exigir que os clientes conectados à rede conheçam detalhes da dispersão geográfica dos membros participantes. Conforme ressalta Johnson, na transmissão *multicast*, somente uma cópia da mensagem é enviada para uma conexão da rede. As cópias da mensagem são replicadas somente quando o caminho divergir em um roteador [JOH 97].

Por outro lado, os sistemas de comunicação e administração de imagens pela rede envolvem um ambiente digital, onde as imagens são adquiridas, transmitidas, analisadas e armazenadas, usando computadores e tecnologias de comunicação. O *throughput*¹ exigido para comunicação de grande quantidade de imagens é extremamente alto. Algumas exigências para tais tipos de transmissão podem ser vistas em RFC 1458 [BRA 93].

Com a crescente utilização das aplicações que exigem transmissão de dados multimídia de forma *multicast*, verifica-se a necessidade de se maximizar o uso dos

¹ *Throughput*: medida de velocidade de transferência de dados através de um sistema de comunicação complexo, ou a velocidade de processamento de dados em um sistema de computador [DIC 98].

recursos da rede, sendo, portanto, necessário avaliar as melhores alternativas de utilização dos mesmos, buscando o menor tempo de resposta e a mínima utilização de largura de banda possível.

Dentre as diversas técnicas de avaliação de desempenho de rede de comunicações existentes, é utilizada neste trabalho a modelagem analítica através da Rede de Filas. Conforme Hammond [HAM 86], a técnica de avaliação de desempenho tem como objetivo desenvolver e estudar modelos matemáticos que façam uma previsão da performance de tais redes em determinadas condições. A vantagem desta técnica é tornar possível analisar a arquitetura da rede sem que a mesma esteja implantada, de forma mais rápida que nas simulações e não tão limitadas quanto os *benchmarking*. Normalmente, vários projetos podem ser modelados e seus desempenhos podem ser avaliados e comparados. Dentre alguns trabalhos realizados nesta linha, pode-se citar o estudo de avaliação de desempenho de Vídeo sob Demanda sobre RSVP e ATM [TRE 99], entre outros.

Sendo o protocolo IP muito difundido nas redes atuais, este trabalho trata da avaliação de desempenho da transmissão *multicast* em redes IP sobre ATM, fazendo uma análise comparativa entre duas técnicas de transmissão *multicast*: com a utilização de servidor *multicast* (MCS) e através de VC *meshes* (sem utilização de servidor *multicast*). O ponto de vista adotado para esta análise é o tempo de resposta no estabelecimento da conexão *multicast*.

II. TRANSMISSÃO IP MULTICAST

A transmissão *multicast* é utilizada quando um nó deseja enviar a mesma informação a vários nós destinos que formam um grupo, identificados por um único endereço de grupo.

Desta forma, em vez de o emissor enviar “n” cópias do mesmo dado a cada destino, é enviada uma única cópia pela rede, sendo que a mesma é replicada apenas quando os caminhos aos destinos finais forem distintos. Isto leva a uma diminuição no tráfego pela rede.

III. ATM – SYNCHRONOUS TRANSFER MODE

Criado para transmitir informações de qualquer espécie que substitua o sistema de redes telefônicas atuais e as redes especializadas como SMDS e DQDB, a tecnologia B-ISDN (*Broadband Integrated Services Digital Network*) integra em um único meio a transmissão de dados, voz e vídeo, oferecendo ampla largura de banda, possibilitando vídeo sob demanda, interconexões de LANs, transporte de dados em altas velocidades, etc. A tecnologia que torna B-ISDN possível chama-se ATM (*Asynchronous Transfer Mode*) [TAN 97].

ATM é baseado na comutação de células de 53 bytes de comprimento, divididos em duas partes: 5 bytes para o *header* e 48 bytes para o *payload*. O *header* é utilizado pelas interfaces UNI (*User-Network Interface*) e NNI (*Network-Network Interface*) para estabelecer conexões, e pelas *switches* para o encaminhamento das células aos seus

destinos apropriados através das conexões virtuais. A parte referente ao *payload* corresponde às informações do usuário.

O *header* da célula possui dois formatos: UNI (*User to Network Interface*) e NNI (*Network to Network Interface*). As interfaces UNI são responsáveis pelo estabelecimento da conexão entre sistemas finais (*hosts*, roteadores, etc.) e *switches* ATM, enquanto as interfaces NNI são responsáveis pela interconexão entre *switches* ATM.

As redes ATM são orientadas à conexão [TAN 96], portanto, antes de iniciar a transmissão de dados entre a fonte e o destino, um circuito virtual deve ser ativado através da rede ATM.

A *switch* ATM é responsável pelo trânsito das células através da rede. [CIS 99].

Um *endpoint* contém um adaptador com interface para a rede ATM. Pode ser uma estação de trabalho (*workstation*), um roteador, uma unidade de serviço digital, *switches* LAN, codificador/decodificador de vídeo (CODCs) [CIS 99].

A tecnologia ATM define um modelo de referência composto pelas seguintes camadas [CIS 99]: a camada física; semelhante à do modelo OSI; a camada ATM combinada com a camada de adaptação ATM (AAL) é responsável por estabelecer conexões e passar as células através da rede ATM utilizando as informações constantes no *header* das células.

A função da camada de adaptação é possibilitar que as células possam ser utilizadas pelas camadas mais altas, buscando empacotar eficientemente as várias espécies de dados da camada superior, tais como datagramas, voz, vídeo, em formatos apropriados para serem enviados ao receptor [PAR 93]. Embora tenham sido definidas várias camadas AAL, vamos nos ater à camada AAL 5, por suportar o protocolo IP.

ATM é a primeira tecnologia capaz de entregar diferentes tipos de tráfego (voz, vídeo e dados) sobre um único mecanismo digital, suportando aplicações multimídia com Qualidade de Serviço [CAR 98]. Os parâmetros de QoS (qualidade de serviços) estão detalhados em [CAB 97]. Quando um *host* conecta-se a uma rede ATM, firma um contrato de tráfego [CAB 97] com a mesma, baseado na qualidade de serviço, especificando o fluxo de tráfego pretendido.

A. Sinalização

Conforme [CIS 98], quando dois dispositivos ATM necessitam estabelecer uma conexão entre si, o primeiro envia uma solicitação de conexão (sinalização) para a *switch* a ele conectado. Este pedido contém o endereço ATM do destino e os parâmetros de QoS exigidos para a conexão. Se a *switch* possuir em sua tabela o endereço do destino e puder acomodar o parâmetro de QoS para a conexão, esta será ativada.

Todas as *switches* ao longo do caminho até o destino final examinam o pedido e enviam para a próxima *switch*, ativando a conexão virtual. Quando o pedido de sinalização chega ao seu destino e este não puder suportar o parâmetro de QoS desejado, é retornada uma mensagem de rejeição à

origem do pedido. Se o pedido puder ser aceito, o destino retorna uma mensagem aceitando a conexão.

ATM Forum especificou UNI 3.0/3.1, que está sendo gradativamente atualizado pela versão UNI 4.0. A especificação do protocolo de sinalização P-NNI 1.0, que interconecta *switches* em redes privadas, pode ser encontrada em [ATM 99].

O protocolo de sinalização se encontra na camada ATM no modelo de referência.

IV. MULTICAST EM ATM

Na transmissão *multicast* em redes ATM, um *host* é conectado a vários pontos finais de forma unidirecional. As replicações das células podem ser feitas tanto nas *switches* ATM quanto pelos próprios *hosts*.

Segundo Alles [ALL 95], nas transmissões *multicast* ATM as conexões são unidirecionais, ou seja, somente o nó raiz (fonte) pode transmitir informações aos nós folhas.

Existem algumas propostas de *switches* ATM que possibilitam a transmissão *multicast*, entre as quais pode-se citar a *switch multicast* ATM Abacus [CHA 97], *Switch Multicast* ATM Clos-Knockout [CHA 98], entre outros.

V. IP SOBRE ATM

Como as tecnologias LANs e WANs atuais existentes baseiam-se em protocolos de camadas de rede como IP, a tecnologia ATM deve suportar conexão com este protocolo.

Segundo Alles [ALL 95], há duas formas de utilizar o protocolo da camada de rede IP sobre a rede ATM: através de LANE (LAN Emulation) e IP Clássico (Classical IP). Conforme descrito em [CAB97], a Emulação de LAN (LANE) faz com que uma rede ATM comporte-se como uma LAN Ethernet ou Token Ring, operando a uma velocidade muito mais rápida.

RFC 1577 de Laubach [LAU 94] define o protocolo IP Clássico sobre ATM, que suporta resolução automática de endereço IP para ATM, introduzindo o conceito de sub-rede IP Lógica, LIS (Logical IP Subnet). LIS consiste de um grupo de nós IP (*hosts*, roteadores), que se conectam a uma única rede ATM e que pertencem à mesma sub rede IP lógica [CISC 95]. Assim, cada LIS deve ser configurada com o endereço que permita acessar um único servidor de resolução de endereços ATMARP. O servidor ATMARP tem como função fazer o mapeamento do endereço IP com o respectivo endereço ATM do *endpoint* (*host*, roteador, etc).

VI. MULTICASTING UTILIZANDO UNI 3.0/3.1

A interface de rede de usuário UNI 3.0/3.1, especificada por ATM Forum para suportar conexão ponto a ponto bidirecional entre dois nós em uma rede ATM, não permite endereçamento *multicast*. Para que esta interface possa suportar transmissão *multicast*, é necessário definir alguns serviços de registro e comportamentos de *endpoints*, dentre eles o servidor de resolução de endereços *multicast* (MARS)[ARM 96].

A. Multicast Address Resolution Server - MARS

MARS é uma extensão do serviço fornecido pelo servidor ATMARP para suportar a transmissão *multicast*. O servidor MARS faz um mapeamento do endereço de grupo do protocolo IP com as interfaces ATM que compõem este grupo.

Os nós (*endpoints*) que desejarem se juntar ou deixar um determinado grupo devem fazê-lo através do servidor de resolução de endereços *multicast* (MARS). MARS mantém os nós informados das alterações ocorridas no grupo. Os *endpoints* são identificados pelos seus endereços ATM *unicast*.

Define-se *cluster* como conjunto de interfaces ATM (*endpoints*) que buscam comunicação *multicast* entre si, registrados a um único MARS. *Cluster* são, portanto, *endpoints* que se registram a um servidor de endereços MARS para atender as necessidades de tráfego *multicast*.

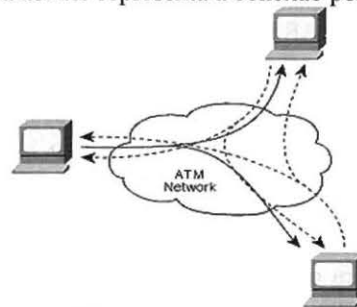
A transmissão *multicast* suportada pelo protocolo UNI 3.0/3.1 apresenta algumas restrições: as conexões virtuais estabelecidas são unidirecionais e somente o nó raiz (fonte) pode adicionar ou remover nós folhas.

A transmissão *multicast* pode ser feita através de:

- VC *meshes*;
- Servidores *multicast*.

A.1 VC Meshes

A figura abaixo representa a conexão por VC *Meshes*.



Fonte: ALLES, 1995, p.10. [ALLE 95].

Fig.1 Conexão Virtual *multicast* utilizando VC *Meshes*

Neste sistema, cada nó, fonte dos pacotes *multicast*, estabelece sua própria conexão virtual (VC) ponto a multiponto para os nós folhas destinos. Assim, os nós podem ser tanto emissores quanto receptores de dados *multicast*. Para um dado grupo, cada nó possui uma VC ponto a multiponto para emissão de pacotes *multicast*, e termina uma VC para cada emissor ativo do grupo para recebimento dos pacotes destinados ao grupo.

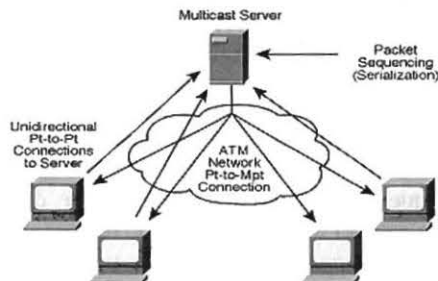
A.2 Servidores *multicast* (MCS – Multicast Servers)

Neste caso, um nó intermediário é escolhido para ser o servidor *multicast*. Cada nó, fonte dos pacotes *multicast*, estabelecerá uma conexão virtual ponto a ponto com este servidor. Estabelecida a conexão, a fonte dos dados *multicast* enviará os pacotes a serem transmitidos ao grupo de forma *unicast* ao MCS. O servidor *multicast*, por sua

vez, estabelecerá e administrará uma conexão virtual ponto a multiponto com todos os nós pertencentes ao grupo, por onde enviará os pacotes de forma *multicast*.

É importante notar que o *host* emissor dos dados *multicast* não necessita saber quais são os membros pertencentes ao grupo, sendo esta a responsabilidade do MCS. O *host* emissor deve apenas conhecer o endereço do MCS.

A fig. 2 mostra a operação de um servidor *multicast*.



Fonte: ALLES, 1995, p.9, [ALL 95].

Fig. 2: Operação do Servidor Multicast – MCS

B. Funcionamento do MARS

Os *endpoints* que desejarem se juntar a um *cluster* devem ser configurados com o endereço ATM do nó onde se localiza o MARS que serve ao *cluster*. O MARS guarda uma tabela de mapeamento contendo {endereço de grupo IP, ATM.1, ATM.2, ... ATM.n}.

O caminho de sinalização inicial, entre um cliente (*endpoint*) e seu MARS, é uma VC ponto a ponto bidirecional. Esta VC estabelecida pelo cliente é utilizada para enviar pedidos e receber respostas do MARS.

Na hipótese de estabelecimento da conexão via VC *Meshes*, quando um pacote com endereço IP *multicast* é recebido pela interface para transmissão, o *endpoint* verifica se já existe um caminho para o grupo *multicast*. Caso ainda não exista, o MARS é requisitado através de uma mensagem para que o mesmo informe o conjunto de *endpoints* ATM que representam o grupo. O MARS responde através de uma sequência de mensagens *MARS_MULTI* por onde é retornado o conjunto de endereços ATM dos membros do grupo.

Se as mensagens de resolução de endereços forem recebidas com sucesso, o *endpoint* pode estabelecer uma conexão ponto a multiponto. O emissor da mensagem *multicast* envia, então, uma mensagem para estabelecer um VC ponto a ponto com o primeiro membro do grupo. A seguir, emite uma solicitação de conexão aos demais endereços {ATM.2, ... , ATM.n}, adicionando estes membros à VC estabelecida. O pacote é, então, enviado sobre esta VC ponto a multiponto.

Na hipótese de se utilizar Servidor Multicast, pode-se utilizar mais de um MCS. Neste caso, o MARS deve guardar dois conjuntos de tabelas de mapeamento de endereços de grupo: um mapa de *hosts* contendo o mapeamento do endereço de grupo IP com os endereços ATM dos *hosts* e roteadores que desejam receber pacotes destinados ao grupo e um mapa de MCSs que mapeia o

endereço de grupo IP com endereços ATM dos MCSs que servem o grupo.

C. Funcionamento do MCS

O servidor *multicast* MCS administra uma única conexão virtual *multicast*. A arquitetura de MCS não suporta mais que um grupo ao mesmo tempo, pois não tem como diferenciar tráfego destinado a diferentes grupos.

D. Avaliação de Desempenho

As questões de desempenho assumem uma grande importância no desenvolvimento das redes de computadores, pois elas permitem avaliar a habilidade de a rede suportar a sua demanda de utilização. Esta demanda de utilização é fator essencial para a especificação, desenvolvimento e uso adequado do sistema [MOU 86]. Estudos de avaliação de desempenho têm sido desenvolvidos como pode ser constatado em [TRE 96], onde se analisa o desempenho de multiprocessador CPER, assim como em [TRE 99] onde se avalia o desempenho do serviço de Vídeo sob Demanda utilizando RSVP em redes ATM.

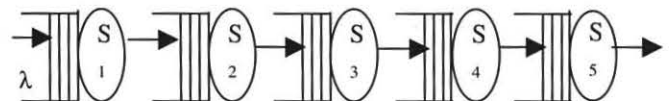
Onvural [ONV 94] relaciona algumas métricas em redes ATM, definindo vários fatores que determinam o atraso na transferência de células em uma rede ATM. De Prycker [PRY95] faz um estudo das características de atraso da rede ATM, dividindo a mesma em diferentes partes, cada qual contribuindo individualmente para o atraso total da rede. Este trabalho utiliza alguns resultados destes autores na formulação do modelo considerado.

VII. MODELO PROPOSTO

Neste trabalho está sendo apresentado um modelo de rede de filas para avaliar o desempenho de serviço *multicast* em redes IP sobre ATM através dos sistemas: VC *Meshes* e com utilização de servidor *multicast*. Este modelo permite avaliar o seu desempenho, concluindo com uma análise comparativa dos resultados obtidos.

A. VC Meshes

O modelo da fig 3.,baixo apresentado, compõe-se de 5



servidores:

Fig. 3: Rede de Filas para VC Meshes.

A.1 Servidor - S₁

Este servidor representa o estabelecimento da conexão entre o *host* emissor dos dados *multicast* e o MARS.

Segundo Onvural [ONV94], esta métrica não é um parâmetro específico de ATM. É principalmente determinado pelo atraso de processamento da mensagem através da rede. A recomendação I.352 da ITU-T define o valor previsto para atraso de ativação de conexão em ISDN de 64 kbps para conexões até 27.500 km. Este mesmo autor assume que o atraso de processamento para ATM é similar aos atrasos em ISDN, sendo que a média de atraso deve ser menor do que 4.500 ms. Neste trabalho está sendo considerado o atraso médio no estabelecimento da conexão por quilometro como sendo 4.500 ms/27.500 km. Desta forma, o atraso total deste servidor será função da distância entre os dispositivos considerados:

$$(1) S_1 = 0,000163636 \times d \text{ s/km}$$

Sendo “d” a distância entre os dois dispositivos

A.2 Servidor – S₂

Este servidor representa o envio da mensagem de solicitação do mapeamento do endereço de grupo (endereço IP *multicast*) para os *hosts* ATM pertencentes ao grupo. Esta mensagem possui aproximadamente 60 bytes.

Segundo De Prycker [PRY95], os parâmetros que contribuem para o atraso total da rede são:

$$(2) S_2 = TD + PD + FD + DD + QD$$

Sendo:

TD = atraso da transmissão da mensagem de solicitação de resolução de endereços

PD = atraso de empacotamento

FD = atraso de comutação

QD = atraso de enfileiramento

DD = atraso de desempacotamento

Estes parâmetros, acima relacionados, estão sendo considerados neste trabalho com base na obra de Onvural [ONV 94].

- Atraso de transmissão - TD

É determinado pela distância e o meio de transmissão independente do conceito de ATM. Considerando que a rede utiliza a tecnologia SONET (Synchronous Optical Network) sob 2 hierarquias básicas de multiplexação OC-3 (taxa bruta de 155,52 Mbps) e OC-12 (taxa bruta de 622,08 Mbps), o tempo de atraso de transmissão do pedido é função do meio utilizado e da distância do *host* emissor ao MARS, uma vez que a largura de banda necessária por esta mensagem é extremamente baixa. Definindo:

TD_{MARS}: o atraso proveniente da solicitação de resolução de endereço feito pelo *host* fonte dos dados *multicast* ao MARS;

TD_i : o atraso gerado em função do meio de comunicação/distância

d : distância entre o *host* emissor e o MARS.

Temos:

$$(3) TD_{MARS} = TD_i \times d$$

Segundo De Prycker, os valores dos atrasos Tdi em função do meio utilizado são descritos na tabela abaixo [PRY 95]:

TABELA I
Atraso de Propagação em Diferentes Meios

Tdi	Atraso de propagação
Cabo coaxial	4 µs / km
Cabo fibra óptica	5 µs / km
Cabo coaxial submarino	6 µs / km
Satellite (14.000 km de altitude)	110 ms / km
Satellite (3.600 km de altitude)	360 ms / km

- Atraso de empacotamento – PD

Refere-se ao preenchimento das células ATM. Segundo De Prycker [PRY 95], este atraso é função do comprimento da célula e da velocidade que a fonte está sendo gerados os bits.

Verificando a proporcionalidade através dos estudos obtidos por Onvural [ONV 94], pode-se concluir que o valor de PD para 48 bytes de *payload* de ATM é aproximadamente de 6000 ms, independente da velocidade de transmissão.

- Atraso de comutação de células – FD

Observada a proporcionalidade para células de 48 bytes de *payload* (ATM) segundo estudos de Onvural[ONV 94], o atraso por *switch* vale em média 24 ms para 150 Mbps e 6 ms para 600 Mbps. O atraso total de comutação de células é função do número de *switches* atravessadas pela mensagem. Considerando “N” o número de *switches* por onde a mensagem deverá atravessar:

$$(4) FD = N \times 0,024 \text{ segundos para 150 Mbps e}$$

$$(5) FD = N \times 0,006 \text{ segundos para 600 Mbps.}$$

- Atraso de desempacotamento e (DD) e atraso de enfileiramento (QD).

O atraso no enfileiramento é função da carga da conexão, da probabilidade de perda de pacotes aceitável e da implementação das *switches*. Considerando carga da rede em torno de 80%, para uma probabilidade de perda de pacotes da ordem de 10⁻¹⁰, o atraso aproximado de enfileiramento e desempacotamento pode ser considerado em média de 75.10⁻⁶ segundos para 150 Mbps e 18,75.10⁻⁶ segundos para 600 Mbps por *switch*. [ONV 94].

Portanto, o atraso QD/DD será:

$$(6) \quad N \times 75.10^{-6} \text{ segundos para 600 Mbps e}$$

$$(7) \quad N \times 18,75.10^{-6} \text{ segundos para 150 Mbps}$$

com N = número de *switches* na rede.

Resumindo, chamando o atraso deste servidor de S_{MARS} temos:

para 150 Mbps:

$$(8) \quad S_2 = S_{MARS} = TDi \cdot d + 6.10^{-3} + 99.10^{-6} \cdot N$$

Para 600 Mbps

$$(9) \quad S_2 = S_{MARS} = TDi \cdot d + 6.10^{-3} + 24,75.10^{-6} \cdot N$$

A.3 Servidor – S_3

Recebido o pedido de resolução de endereços, o MARS faz o mapeamento do endereço IP *multicast* para os respectivos endereços *unicast* ATM dos *host* pertencentes ao grupo, retornando estes endereços ao *host* solicitante, através de uma sequência de mensagens “*Mars_Multi*”.

Considerando negligível o tempo de processamento de resolução de endereço pelo MARS, o servidor 3 representará apenas o tempo de transmissão da mensagem ou das mensagens “*Mars_Multi*”, contendo a sequência de endereços *unicast* ATM, do MARS ao *host* solicitante.

Assim, havendo o mapeamento de endereços, o MARS retorna uma ou mais mensagens “*Mars_Multi*” transportando os múltiplos endereços *unicast* ATM {ATM.1, ATM.2,..., ATM.n}. O limite do tamanho de cada mensagem “*Mars_Multi*” é o valor da MTU (máxima unidade de transmissão) permitida para o transporte de mensagem ATM. Por *default*, o valor da MTU para ATM é fixado em 9.180 bytes, similar ao SMDS. Assim, é possível transportar 456 membros em uma única mensagem.

Para grupos com algumas centenas de membros de grupo, a largura de banda utilizada por esta mensagem é desprezível. Desta forma, o atraso provocado por este servidor é similar ao atraso gerado pelo servidor 2.

A.4 Servidor – S_4

O servidor 4 representa o atraso referente ao estabelecimento da conexão do *host* emissor com o primeiro *host* (membro do grupo).

Similarmente ao servidor 1 (S_1), o atraso considerado é de 0.000163636 s / km

A.5 Servidor – S_5

O *host* emissor estabelece conexão com os demais *hosts*. Neste caso, admite-se que os estabelecimentos das

conexões sejam feitos em paralelo e o atraso gerado por este servidor é igual ao servidor 4 anterior.

O modelo representando o tempo de serviço de cada servidor fica, então, assim resumido:

$$(10) \quad S_1 = S_4 = S_5 = 0.000163636 \times d \text{ seg / visita.}$$

Chamando V_i e λ_i , a taxa de visitas em um determinado servidor “*i*” (ou seja, o número de vezes que a mesma mensagem atravessa um determinado servidor) e a taxa de chegadas no servidor “*i*” respectivamente, tem-se:

$$(11) \quad V_1 = V_2 = V_3 = V_4 = V_5 = 1 \text{ visita,}$$

ou seja, cada mensagem visita apenas uma vez cada servidor.

Sendo λ = taxa média de chegada de pedidos de emissão de dados *multicast*, tem-se que: $\lambda_i = V_i \times \lambda = \lambda$.

Pelo Teorema da Taxa de Processamento,

$$(12) \quad U_i = \lambda \cdot S_i \text{ sendo } U_i = \text{fator de utilização e} \\ S_i = \text{tempo médio de serviço}$$

Associando ao Teorema de Tempo de Resposta:

$$(13) \quad R_i = \frac{S_i}{1 - U_i} = \frac{S_i}{1 - \lambda \cdot S_i}$$

Conclui-se, portanto,

$$(14) \quad R_1 = R_4 = R_5 = R_{\text{conexão}} = \frac{0.000163636 d}{(1 - \lambda \cdot 0.000163636 d)}$$

$$(15) \quad R_2 = R_3 = \frac{S_{MARS}}{(1 - \lambda \cdot S_{MARS})}$$

Pelo Teorema de Burke para redes Abertas:

$$(16) \quad R = \sum V_i \cdot R_i = \sum R_i$$

Das equações (8), (9), (14) e (16) deduz-se o tempo necessário para o estabelecimento de conexão *multicast* para VC *Meshes*:

$$(17) \quad R = 3 \cdot R_{\text{conexão}} + \frac{2 \cdot S_{MARS}}{(1 - \lambda \cdot S_{MARS})}$$

B. Modelo de Filas com Servidor Multicast

O modelo apresentado na figura 4 a seguir compõe-se de 7 servidores descritos a seguir:

B.1 Servidor 1- S_1^{MCS}

O *host* emissor estabelece conexão com o MARS

B.2 Servidor 2 - S_2^{MCS} :

O *host* emissor envia ao MARS a solicitação de resolução de endereços IP *multicast*;



Fig.4 Rede de Filas com Servidor *Multicast* (MCS)

B.3 Servidor 3 - S^{MCS}_3 :

O MARS envia a mensagem, contendo o endereço do servidor *multicast* (MCS), ao *host* emissor.

B.4 Servidor 4- S^{MCS}_4 :

O *host* emissor estabelece conexão com o MCS;

B.5 Servidor 5- S^{MCS}_5 :

O MCS estabelece conexão com o primeiro endereço do grupo *multicast*;

B.6 Servidor 6 - S^{MCS}_6 :

O MCS estabelece conexão com os demais membros do grupo *multicast*;

B.7 Servidor 7 - S^{MCS}_7 :

O *host* emissor envia os dados *multicast* ao MCS

Neste modelo, considera-se que o servidor *multicast* MCS já possui o mapeamento de endereços {IP *multicast*, ATM.1, ATM.2,...ATM.n} em seu *cache*.

Comparando os parâmetros desta rede de filas à rede anterior (VC *Meshes*), tem-se:

O servidor S^{MCS}_1 equivale ao servidor S_1 do modelo anterior = 0,000163636 d segundos eq.(1);

O servidor S^{MCS}_2 equivale ao servidor S_2 do modelo anterior = S_{MARS} eq. (8) e (9);

O servidor S^{MCS}_3 possui o mesmo tamanho da mensagem do servidor anterior = S_{MARS} eq (8) e (9);

O servidor S^{MCS}_4 equivale ao primeiro servidor = 0,000163636 d segundos eq. (1);

O servidor S^{MCS}_5 equivale ao servidor S_4 do modelo anterior = 0,000163636 d segundos eq. (1);

O servidor S^{MCS}_6 equivale ao servidor S_5 do modelo anterior = 0,000163636 d segundos eq.(1);

O servidor S^{MCS}_7 representa o tempo gasto para enviar os dados *multicast* de forma *unicast* do *host* emissor ao MCS.

O modelo de filas ficará então resumido a:

$$(18) S^{MCS}_1 = S^{MCS}_4 = S^{MCS}_5 = S^{MCS}_6 = 0,000163636 \text{ seg / visita.}$$

Sendo $V_1 = V_2 = V_3 = V_4 = V_5 = V_6 = V_7 = 1$ visita, e λ , a taxa média de chegada de pedidos de emissão de dados *multicast*, tem-se que $\lambda_i = \lambda$

Pelo Teorema da Taxa de Processamento associado com o Teorema de Tempo de Resposta mencionado anteriormente:

$$(19) R_1 = R_4 = R_5 = R_6 = R_{\text{conexão}} = \frac{0,000163636 \text{ d}}{(1 - \lambda \cdot 0,000163636 \text{ d})}$$

$$(20) R_2 = R_3 = \frac{S_{MARS}}{(1 - \lambda S_{MARS})}$$

Das equações (8), (9), (13), (16), (19) (20) deduz-se o tempo necessário para o estabelecimento de conexão *multicast* com utilização de servidor *multicast* (MCS):

$$(21) R = 4 \cdot R_{\text{conexão}} + \frac{2 \cdot S_{MARS}}{(1 - \lambda \cdot S_{MARS})} + \frac{S_{\text{APLICATIVO}}}{(1 - \lambda \cdot S_{\text{APLICATIVO}})}$$

VIII. RESULTADOS, DISCUSSÃO E CONCLUSÕES

Analisando um estudo de caso onde admite-se como meio de transmissão a fibra ótica, distância média de 1.000 km entre os membros do grupo, média de 10 *switches* entre a fonte emissora e os destinos, taxa média de chegada de pedidos de transmissão *multicast* de 0.0002 pedidos/segundo: nesta situação, o tempo de resposta no sistema VC *Meshes* para o estabelecimento da conexão *multicast* é de aproximadamente 23 milissegundos. Comparando as equações (17) e (21), conclui-se que, em relação ao tempo de estabelecimento de conexão *multicast*, a diferença entre o sistema com servidor *multicast* e o sistema VC *Meshes* resume-se ao tempo de estabelecimento da conexão e a emissão dos dados *multicast* entre o *host* emissor de dados *multicast* e o servidor *multicast* MCS. Se a transmissão dos dados *multicast* variar entre 10 a 40 minutos, por exemplo, este tempo, basicamente, representará a diferença entre os dois sistemas, uma vez que é bem superior aos 23 milissegundos obtidos em VC *Meshes*.

Comparando-se agora apenas no sistema VC *Meshes*, variando-se a taxa de transmissão de 150 Mbps e 622 Mbps em função da distância média entre os membros do grupo (considerando uma média de 10 *switches* ATM a cada 1.000 km), conclui-se que para pequenas distâncias entre os *hosts*, a diferença entre os tempos de estabelecimento de conexão *multicast* é desprezível. Para distâncias acima de 5.000 km, o tempo de estabelecimento de conexão para 622 Mbps torna-se aproximadamente 88% do tempo gasto para conexão em 150 Mbps.

Uma das vantagens de se utilizar o servidor *multicast* está justamente na economia dos recursos da rede. Isto porque, quando se utiliza VC *Meshes*, todos os *hosts* emissores de dados *multicast* devem estabelecer conexões com todos os *hosts* membros do grupo, enquanto que,

quando se utiliza servidores *multicast* (MCS), cada *host* estabelece apenas 2 conexões: uma para enviar dados ao MCS e outra para receber dados do MCS.

Então, para ilustrar, supondo uma aplicação de videoconferência com sistema de distribuição de vídeo necessitando de uma taxa de bits de 1.484 kbps, conforme Onvural [ONV 94] tem-se que, para rede de 150 Mbps, suportaria apenas um grupo composto de 10 membros em sistema de VC *Meshes*, onde todos os *hosts* devem estabelecer conexões com todos os membros do grupo. Assim, 10 membros x 10 conexões/cada x 1,484 Mbps = 148,4 Mbps de consumo de recursos da rede. Para redes que suportam 622 Mbps, comportariam 20 membros em VC *Meshes*.

Utilizando servidor *multicast*, a rede comportaria um grupo de 50 membros, ou 5 grupos de 10 membros. Neste último caso, seria interessante ter 5 servidores *multicast*, para que a transmissão de um grupo não seja interferido pela transmissão de outro, evitando o gargalo no MCS.

O modelo apresentado neste trabalho encontra-se em fase de testes e validação através de simulação de casos configurados.

Como principais contribuições, este modelo permite de forma bastante rápida e simples avaliar-se e/ou comparar-se o desempenho das duas técnicas de transmissão IP *multicast* sobre ATM: com servidor *multicast* e com VC *Meshes*, bem como obter medidas para projeto, desenvolvimento ou operação desses sistemas.

AGRADECIMENTOS

Nossos agradecimentos ao apoio obtido pelo programa de Pós-Graduação em Ciência da Computação da UFSCar e da Missão Salesiana do Mato Grosso de Lins.

REFERÊNCIAS

- [ALL 95] ALLES, Anthony. *ATM Internetworking*. White Paper, Cisco System, 1995. URL: <http://www.cisco.com/warp/public/614/html>. Consultado em 18/02/00.
- [ARM 96] ARMITAGE, G., BELLCORE. *Support for Multicast over UNI 3.0/3.1 based ATM Networks*. RFC 2022, 1996.
- [AMS 92] AMSTRONG, S., FREIER, A., MARZULLO, K. *Multicasting Transport Protocol*. RFC 1301, 1992.
- [ATM 99] The ATM Forum – The Technical Committee. *PNNI Transported Address Stack version 1.0*. 1999. AF-CS-0115.000.
- [BRA 93] BRAUDES, R., ZABELE, S. *Requirements for Multicast Protocol*. RFC 1458, 1993.
- [CAB 97] CABLETRON Systems. *ATM Tecnology Guide*. 1997.
- [CAR 98] CARRIEDO, Maria Isabel G. *ATM Origins and State of the Art*. Barcelona, 1998. Consultado em 06/09/2000.
- [CIS 98] CISCO Systems. *Asynchronous Transfer Mode*. URL: <http://www.ics.muni.cz/cisco/data/doc/cintrnet/ito/55755.htm>, 1998. available: 08/06/2000.
- [CIS 99] CISCO Systems. *Asynchronous Transfer Mode*. White Paper, URL: http://www.cisco.com/univercd/cc/td/doc/cisint/wk/ito_doc/atm.htm, 1999. available: 10/02/2000.
- [CHA 98] CHAN, King-Sun. et.al. *Clos-Knockout: A Large-Scale Modular Multicast Atm Switching*. IEICE Transaction Communication, v. E81B, n.2., p.266-275, 1998.
- [CHA 97] CHAO, H.Jonathan. et.al. *Design and Implementation of Abacus Switch: A Scalable Multicast ATM Switch*. IEEE Journal on Selected Areas in Communications, v.15, n. 5, p.830-843, 1997.
- [DIC 98] *DICIONÁRIO de Informática*. Microsoft Press, 3. ed. Rio de Janeiro : Campus, 1998.
- [HAM 86] HAMMOND, Josheph L., O'REILLY, Peter J.P. *Performance Analysis of Local Computer Networks*. Massachusetts: Addison-Wesley, 1986.
- [JOH 97] VICKI, J., JOHNSON, M. *How IP Multicast Works*. Stardust.com, Inc., 1997. URL: <http://www.ipmulticast.com/community/whitepapers/howipmcworks.html>. Consultado em 27/03/2000
- [KOS 98] KOSIUR, D. *IP Multicasting*. John Wiley & Sons, Inc., 1998.
- [LAU 94] LAUBACH, M., HALPERN, J. *Classical IP and ARP over ATM*. RFC 1577, 1998.
- [MOU 86] MOURA, J. A.B. et. al. *Redes Locais de Computadores: Protocolos de Alto Nível e Avaliação de Desempenho*. McGraw-Hill, São Paulo: 1986.
- [ONV 94] ONVURAL, R.O. *Asynchronous transfer mode networks: performance issues*. Boston: 1994
- [PAR 93] PARTRIDGE, C. *Gigabit Networking*. Addison-Wesley, 1993
- [PRY 95] PRYCKER, Martín de. *Asynchronous Transfer Mode: solution for broadband ISDN*. Prentice Hall, London: 1995.
- [TAN 96] TANENBAUM, Andrew S. *Computer Networks*. 3 ed., Prentice Hall, 1996.
- [TRE 96] TREVELIN, L. C., FERNANDES, M. M. *Analytical performance Modelling of the Cper Multiprocessor*. In: Parallel And Distributed Processing Techniques And Applications, (PDPTA'96) International Conference, Sunnyvale, California: 1996.
- [TRE 99] TREVELIN, Luis C., MESQUITA, T.M.S., *Performance Modeling of Video-On-Demand Service Over RSVP and ATM*. International Conference on Applied Informatics AI'99, Innsbruck, Áustria, fev.1999.