

Administração de Sistemas (ASIST)

Qualidade de Serviço (QoS)

Dezembro de 2014

QoS (“Quality Of Service”)

A qualidade de serviço em rede pode ser expressa por diversos parâmetros mesuráveis (parâmetros de QoS), por exemplo:

- débito de dados (“throughput”)
- taxa de erros e perda de dados
- atraso na transmissão (“latency”)
- variação do atraso (“jitter”)

Oferecer garantias de que determinados valores destes parâmetros não são ultrapassados para os serviços mais importantes exige um tratamento preferencial de algum tráfego em detrimento de outro.

Existem duas abordagens ligeiramente diferentes:

- **HARD QoS** (“guaranteed service”) – determinados recursos ficam reservados (tipicamente uma percentagem da largura de banda), mesmo que esse recurso não esteja a ser usado está sempre indisponível para outros tipos de tráfego.
- **SOFT QoS** (“differentiated service”) – neste caso não há reserva de recursos, os recursos são partilhados pelos vários tipos de tráfego, embora dando preferência ao tráfego mais importante.

SOFT QoS – “Differentiated services” (DiffServ)

O “Soft QoS” caracteriza-se pela aplicação de tratamentos diferenciados aos diferentes tipos de tráfego sem assumir quaisquer compromissos mesuráveis.

Todos os recursos de rede estão sempre disponíveis para qualquer tipo de tráfego, garante por isso sempre uma taxa de utilização dos recursos de 100%. Nunca ocorrem situações em que existe tráfego retido num “router” apesar de a ligação de saída não estar ocupada a 100%.

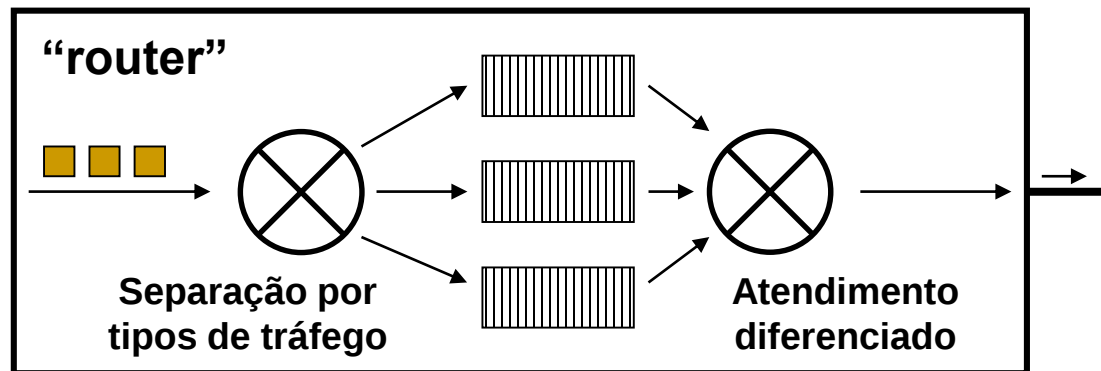
A desvantagem do “Soft QoS” é que não dá garantias absolutas. Se pretender que determinado tipo de tráfego tenha sempre disponível para si determinados recursos de rede, o “Soft QoS” não é a solução.

O “Soft QoS” apenas garante que determinados tipos de tráfego são mais favorecido que outros através de um tratamento diferenciado. Não garante valores absolutos de parâmetros QoS para nenhum tipo de tráfego.

SOFT QoS – gestão de filas

As filas de saída dos nós intermédios são o ponto de estrangulamento nas comunicações, em caso de congestionamento provocam atrasos na transmissão, redução do débito de dados e eventualmente perdas de pacotes.

Para garantir melhores parâmetros QoS para o tráfego mais importante são criadas várias filas de saída paralelas de acordo com o tipo de tráfego. O tráfego de entrada é separado em várias filas que vão ser atendidas de forma diferenciada.

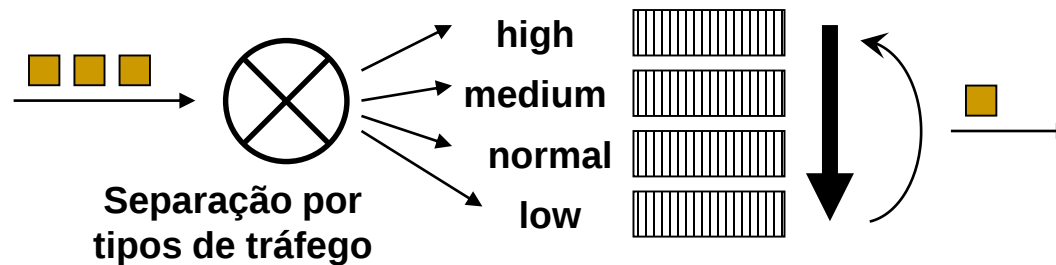


A classificação do tipo de tráfego pode ser realizada localmente no nó intermédio ou pode já ter sido realizada noutro nó, estando nesse caso o pacote marcado com um valor apropriado nos bits de precedência.

“Priority Queuing” (PQ)

A gestão por filas prioritárias dá prioridade absoluta ao tráfego mais importante, normalmente são definidas 4 filas com prioridades decrescentes (“high”, “medium”, “normal” e low”).

Sempre que há oportunidade para emitir um pacote (ligação de saída disponível) o algoritmo percorre as filas no sentido de prioridade decendente até encontrar um pacote emitindo-o.



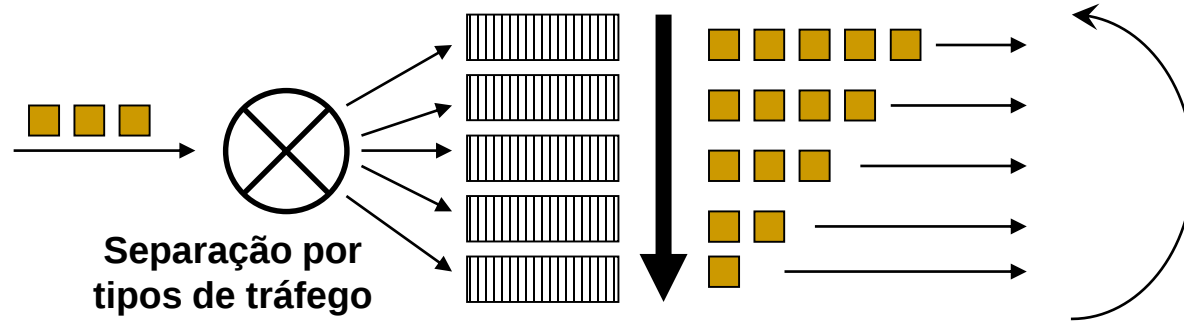
O procedimento repete-se sempre que há oportunidade para emitir um pacote. Devido a este modo de funcionamento um pacote posicionado na frente de uma dada fila só será emitido quando todas as filas superiores estiverem vazias.

Cada fila tem um tamanho fixo, por isso quando existe um volume elevado de tráfego prioritário o restante tráfego fica retido durante muito tempo, podendo mesmo encher a respetiva fila e ser descartado.

“Custom Queuing” (CQ)

O “Custom Queuing” garante que todos os tipos de tráfego são sempre emitidos, contudo o volume ou número de pacotes emitido varia.

A algoritmo percorre as várias filas em ciclo (“round-robin”) e em cada uma das filas processa um volume de informação ou número de pacotes que depende da prioridade dessa fila.



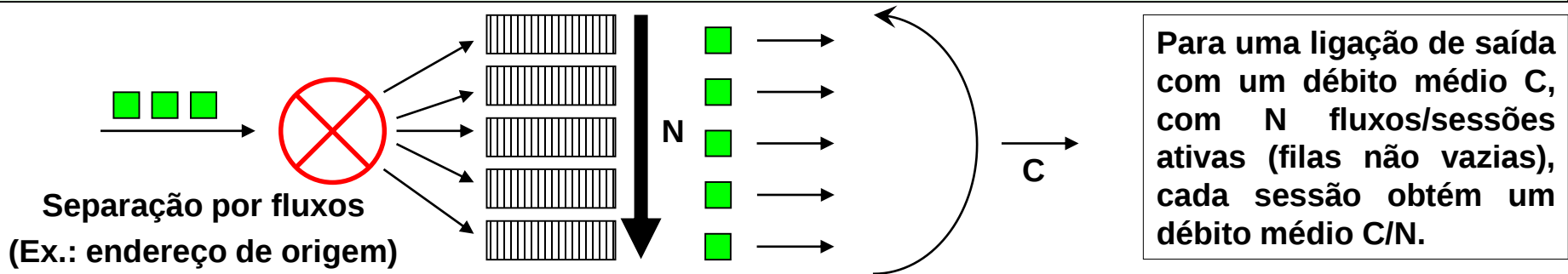
Garante por isso que mesmo o tráfego de prioridade mais baixa tem oportunidade de passar, mas oferece menos garantias para o tráfego de elevada prioridade.

O número de filas existentes e o volume de informação ou número de pacotes a processar por iteração em cada fila podem ser ajustados para obter os resultados desejados, mas trata-se de um ajuste estático.

“Fair Queuing” (FQ)

O “Fair Queuing” (FQ) é uma técnica de distribuição equitativa da largura de banda disponível por um conjunto de sessões identificadas. O nó intermédio começa por identificar fluxos/sessões de informação especialmente através do endereço de origem, mas também destino, protocolos e números de porto.

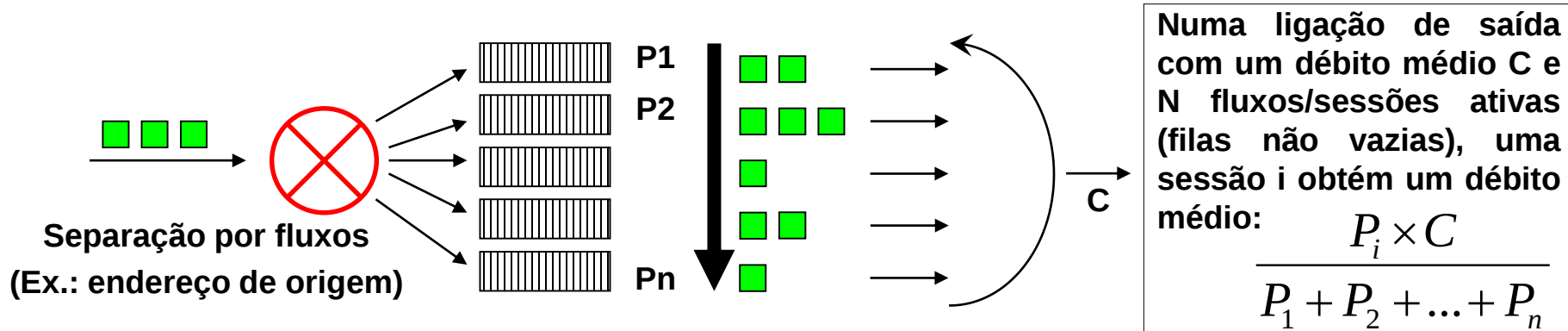
Para cada fluxo de informação identificado é criada uma fila. Processa os pacotes em cada fila de modo a que a capacidade da ligação de saída seja dividida de forma igual por todos os fluxos/sessões existentes.



O “Fair Queuing” permite uma distribuição equitativa da capacidade da ligação de saída pelos vários utilizadores, evitando que uma utilização intensa por parte de alguns prejudique os restantes. É muito usada pelos ISP para garantir oportunidades iguais de utilização aos subscritores independentemente do tráfego que estão a produzir.

“Weighted Fair Queuing” (WFQ)

O “Weighted Fair Queuing” (WFQ) actua de forma semelhante ao FQ, mas a cada fluxo (i) é atribuído um peso diferente (P_i).



Os pesos atribuídos a cada fila podem atender a vários critérios, tipicamente os fluxos são classificados em alto e baixo débito de acordo com o volume de informação que circula. Os fluxos de baixo débito têm prioridade (maior peso).

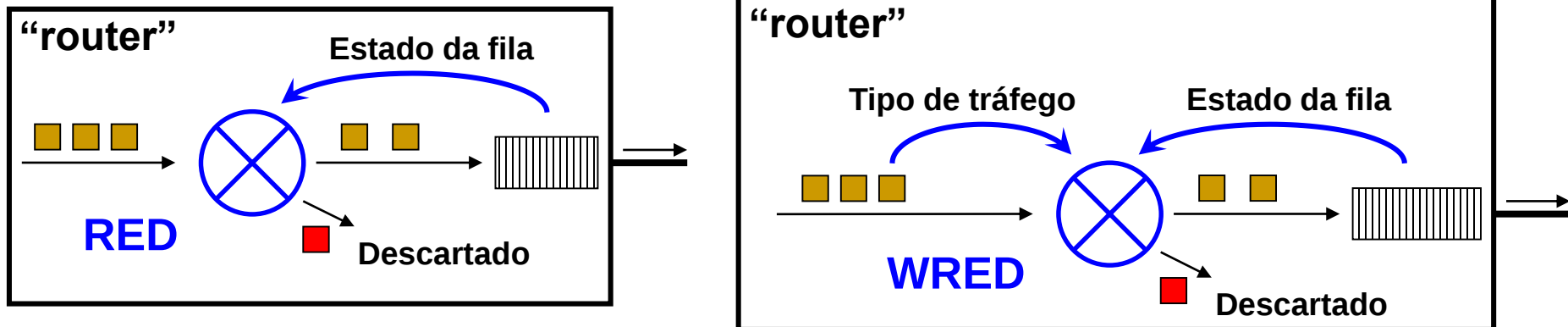
Outros critérios para atribuição de pesos podem ser a classificação direta do tipo de tráfego ou a análise dos bits de precedência IP.

O WFQ dá prioridade ao tráfego interativo e distribui dinamicamente, de acordo com pesos estabelecidos, a largura banda que fica disponível pelo restante tráfego. Desta forma o WFQ adapta-se automaticamente às alterações no tráfego.

“Weighted RED” – WRED

O RED – “Random Early Detection” permite aos nós intermédios evitar que todos os emissores TCP entrem simultaneamente em “slow start” quando a fila enche. Para o conseguirem quando a fila cresce acima de determinado valor começam a descartar aleatoriamente pacotes com uma intensidade crescente com o aumento da fila.

O WRED introduz um critério não aleatório para eleição dos pacotes que devem ser descartados primeiro, ou seja os de baixa prioridade.



Normalmente a classificação do tráfego não é realizada no nó intermédio que implementa o WRED, tipicamente o WRED usa os bits de precedência IP com que os pacotes já foram marcados num outro local.

Hard QoS

O “Hard QoS” utiliza uma abordagem de reserva de recursos, tipicamente uma fração da capacidade de transmissão total (vulgarmente designada largura de banda). Sob o ponto de vista de vantagens e desvantagens é o inverso do “Soft QoS”:

Permite reservar estaticamente uma parte dos recursos disponíveis para determinado tipo de tráfego. Há garantias de que os recursos reservados estão sempre disponíveis permitindo definir valores garantidos para parâmetros QoS.

A desvantagem do “Hard QoS” é que não garante uma taxa de utilização dos recursos a 100%. Podem ocorrer situações em que determinado tráfego é retido ou mesmo descartado num “router” apesar de a ligação de saída não estar a ser utilizada a 100%.

No “Hard QoS”, quando é detetado determinado tipo de tráfego em excesso relativamente ao estabelecido para esse tráfego podem ser desencadeadas dois tipos de medidas: “Traffic Shaping” ou “Traffic Policing”.

Hard QoS - “Traffic Policing” e “Traffic Shaping”

O “Traffic Policing” consiste na simples eliminação de todos os pacotes correspondentes a tráfego em excesso, qualquer pico de tráfego pontual será eliminado pelo que tem um efeito nefasto.

O “Traffic Shaping” consiste na retenção dos pacotes correspondentes a tráfego em excesso, os picos de tráfego pontuais levam a um armazenamento dos pacotes, que serão mais tarde transmitidos quando o volume de tráfego for mais reduzido. Compromete o atraso na rede, mas evita que os pacotes sejam eliminados.

Tipicamente os dois mecanismos são usados em conjunto, o prestador do serviço aplica “Traffic Policing” para garantir que o cliente não viola o contrato, o cliente aplica “Traffic Shaping” para garantir que o prestador do serviço não vai eliminar os seus pacotes por aplicação de “Traffic Policing”.

Os dois mecanismos também podem ser aplicados de forma combinada, por exemplo para excessos menores aplicar “Traffic Shaping” e para excessos maiores aplicar “Traffic Policing”. Note-se que aplicando apenas “Traffic Shaping”, quando a fila encher os novos pacotes serão eliminados.

“Committed Access Rate” (CAR)

O CAR é uma técnica “HARD QoS” de “Traffic Policing” em que se definem limites de taxas de transmissão, tipicamente são definidos limites diferentes para diferentes tipos de tráfego. O CAR pode ser imposto no tráfego de entrada dos “routers” ou no tráfego de saída.

O CAR define valores limite para as taxas de transmissão de cada tipo de tráfego, correspondentes a fatias da taxa de transmissão da ligação existente.

Exemplo numa ligação a 100.000.000 bps (100 Mbps):

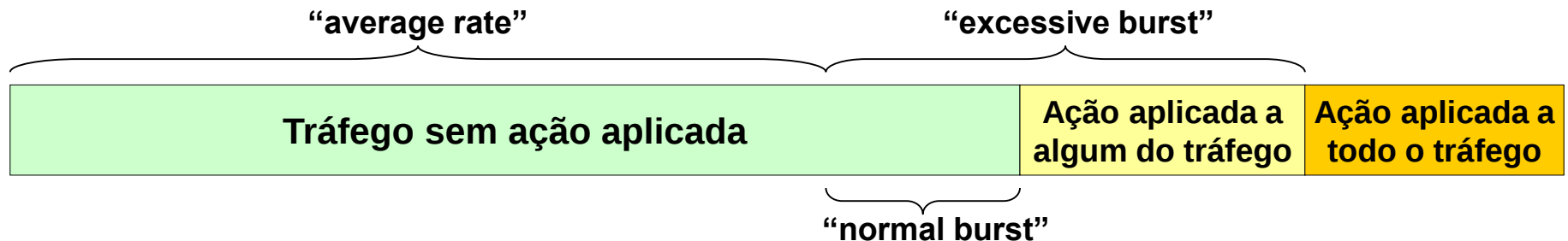
- tráfego não classificado: 50.000.000 bps, ou seja 50%.
- tráfego RTP: 10.000.000 bps
- tráfego HTTP: 30.000.000 bps
- tráfego FTP: 10.000.000 bps

O CAR oferece garantias totais de disponibilidade dentro dos limites de tráfego estabelecidos, tem a desvantagem de ser totalmente estático, no exemplo anterior o tráfego “não classificado” nunca poderá exceder os 50 Mbs, mesmo que não exista nenhum tráfego RTP / HTTP / FTP.

CAR – limites de tráfego

O CAR define os limites de taxa de transmissão através de 3 parâmetros:

- taxa média (“average rate”) – para todo o tráfego abaixo deste limite não é desencadeada qualquer ação – “Committed Information Rate” (CIR)
- pico normal (“normal burst”) – valor a partir do qual começam a ser desencadeadas ações para algum do tráfego – “Committed Burst Size” (CBS)
- pico excessivo (“excessive burst”) - valor a partir do qual começam a ser desencadeadas ações para todo o tráfego – “Excess Burst Size” (EBS)



Na zona compreendida entre o “normal burst” e o “excessive burst” a quantidade de tráfego ao qual é aplicada a ação vai aumentando progressivamente.

CAR – classificação e ações

O objectivo do CAR é garantir determinados valores de taxa de transmissão para determinados tipos de tráfego. A classificação do tráfego no CAR pode ser feita de várias formas, por exemplo com uma ACL (endereços, protocolos, números de porto) ou através dos bits de precedência IP (tráfego já classificado noutra local).

Após a classificação o tráfego é processado no contexto dos limites de taxa de transmissão para ele definidos, nesse contexto esse tráfego poderá ser considerado dentro ou fora dos limites, para cada caso pode ser definida uma ação, por exemplo:

- transmitir
- eliminar (“Traffic Policing”)
- alterar bits de precedência IP e transmitir (aplicação indireta de “Traffic Shaping”)

Tipicamente para o tráfego considerado dentro dos limites a ação será “transmitir” e para o tráfego fora dos limites será “descartar”.

A alteração dos bits de precedência IP pode ser aplicada para alterar a prioridade com que o tráfego vai ser tratado em outros locais da rede, por exemplo em filas ou no contexto do WRED.

Hard QoS – “Integrated services” (IntServ)

Sendo o objetivo do “Hard QoS” garantir valores concretos de QoS para determinados tipos de tráfego, não é suficiente a sua aplicação num único nó intermédio.

Para conseguir garantir valores concretos de QoS para um fluxo de informação entre dois nós finais o “Hard QoS” tem de ser aplicado de forma integrada ao longo de todos os nós intermédios por onde esse fluxo passa.

A arquitetura “IntServ”, muitas vezes considerada oposta à DiffServ” propõe mecanismos que permitem aos nós finais reservar recursos (“Hard QoS”) nos nós intermédios ao longo do percurso do fluxo de informação.

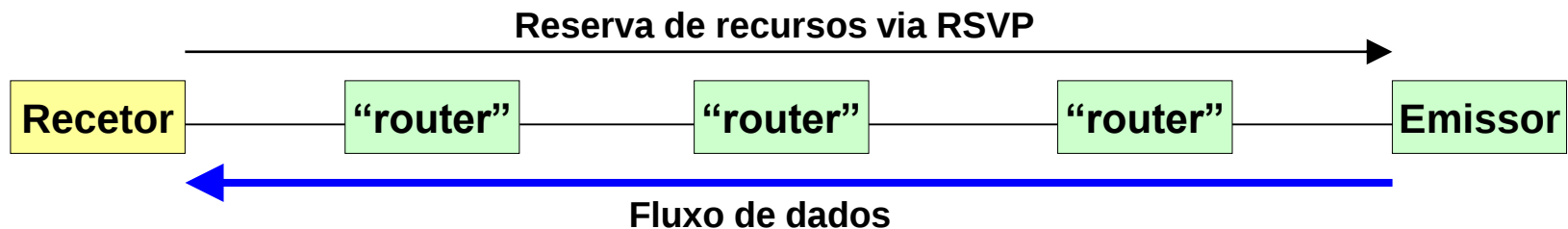
O RSVP (“Resource ReSerVation Protocol”) implementa a arquitetura “IntServ”.

RSVP (“Resource Reservation Protocol”)

O RSVP permite a um nó da rede especificar características QoS que pretende para um determinado fluxo de dados que está a receber. Para esse efeito o protocolo RSVP é usado para comunicar com os agentes RSVP existentes nos vários nós intermédios (“routers”) que configuram a prioridade com que o tráfego correspondente vai ser atendido.

Sob o ponto de vista do RSVP um fluxo é uma sequência de pacotes com o mesmo endereço de destino (nó e opcionalmente número de porto) usando um determinado protocolo de transporte. No entanto o RSVP suporta endereços “multicast”, por isso o endereço de destino pode corresponder a vários nós.

Cada fluxo apenas representam um sentido de circulação do tráfego (“simplex”), cabe ao recetor desse fluxo usar o RSVP para reservar recursos.



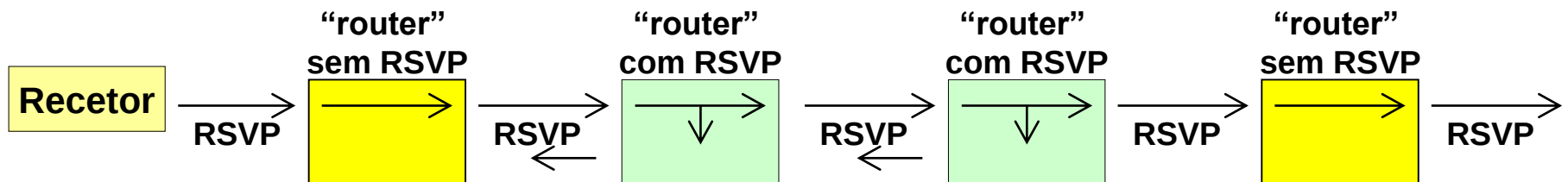
RSVP – Pedidos de reserva

Os pedidos de reserva de recursos devem ser continuamente emitidos pelo nó recetor (normalmente de 30 em 30 segundos), caso contrário as reservas efectuadas nos nós intermédios expiram. A reserva de recurso é por isso dinâmica, ajustando-se às eventuais alterações de encaminhamento na rede.

Os pedidos de reserva de recursos são emitidos com o endereço de destino correspondente ao emissor, seguem por isso o caminho inverso do fluxo de dados a que se vão aplicar. O endereço de destino pode também ser um endereço multicast.

Os nós intermédios que não possuem agente RSVP ignoram os pedidos RSVP efetuando o seu encaminhamento juntamente com o restante tráfego.

Os nós intermédios com agente RSVP interceptam os pedidos RSVP, mas também efetuam o seu encaminhamento juntamente com o restante tráfego para que sejam propagados pelos vários nós intermédios até ao destino. Os pedidos são passados ao agente RSVP que os processa e envia uma resposta ao recetor.



RSVP – propagação dos pedidos de reserva

Os pedidos de reserva de recursos do RSVP são propagados pela rede tendo como origem o recetor do fluxo de dados, os nós intermédios sem RSVP ignoram os pedidos. Os nós com RSVP processam os pedidos e informam o recetor, se o pedido de reserva tiver sucesso o nó intermédio encaminha o pedido para o nó seguinte, mas pode reformular esse pedido alterando parâmetros.

O RSVP agrega pedidos de reserva compatíveis, por essa razão num dado ponto do percurso o pedido pode ser alterado ou mesmo eliminado por corresponder a uma reserva já existente. Isto acontece frequentemente quando são usados endereços multicast.

Quando o pedido consegue completar com sucesso todo o percurso até ao emissor ou quando é eliminado por ser agregado a uma reserva já existente pode ser enviada ao recetor uma mensagem de sucesso. Devido à agregação de reservas, esta mensagem de sucesso dá ao recetor apenas garantias parciais de que foram estabelecidas as condições pretendidas.

RSVP – “flowspec” e “filter spec”

Os pedidos de reserva de recursos do RSVP contêm dois elementos base que são usados pelos agentes nos nós intermédios:

“flowspec” – define o nível QoS pretendido, contém uma classe de serviço RSVP e parâmetros adicionais (taxa mínima, atraso máximo, etc) que são independentes do RSVP. Esta informação é usada para configurar o mecanismo de tratamento diferenciado como por exemplo as filas de saída.

“filter spec” – define a que tráfego se aplica, é usado para configurar o mecanismo de classificação do tráfego, por exemplo através de uma ACL.

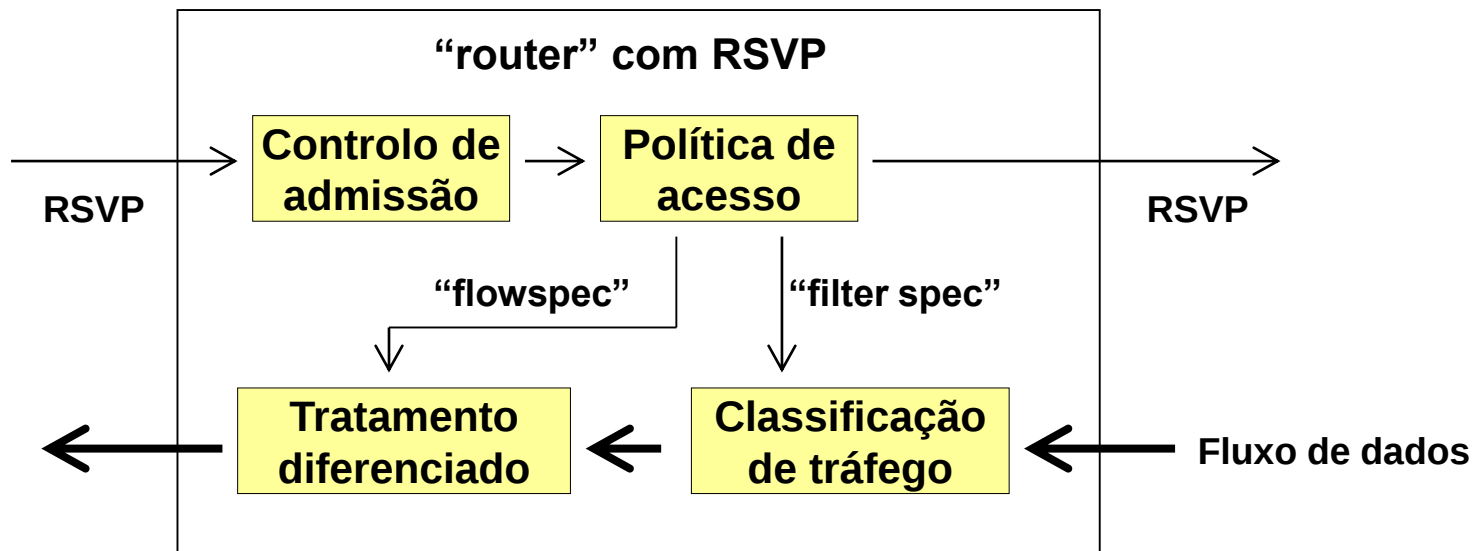
Cada fluxo é classificado numa classe de serviço RSVP:

- “Best-effort” – menos sensíveis às taxas de transmissão e atrasos, exigentes quanto aos erros de transmissão, normalmente exigem a correção automática de erros, tipicamente pelo protocolo TCP.
- “Rate-sensitive” – maior preocupação na manutenção de uma taxa de transmissão garantida e constante.
- “Delay-sensitive” – maior preocupação na manutenção de tempos de atraso na rede garantidos e constantes.

RSVP – Processamento de pedidos

Quando um pedido RSVP é recebido por um nó intermédio, o “controlo de admissão” começa por verificar se existem recursos disponíveis para atender o pedido. Em caso de sucesso, segue-se a verificação da política de acesso que se destina a verificar se o utilizador tem permissão para efectuar a reserva.

Em caso de falha em qualquer dos dois passos anteriores o processo de reserva no nó falha e o recetor é informado com uma mensagem de erros apropriada.



RSVP – estilo de reserva (“RSVP Reservation Style”)

Os pedidos de reserva de recursos do RSVP incluem parâmetros que definem o estilo de reserva como parte do “filter spec”.

Existem duas classes de reserva:

“distinct” – o pedido associa um único emissor ao fluxo

“shared” – o pedido associa vários emissores ao fluxo

Estilos de reserva RSVP

FF (“fixed-filter”) – reserva da classe “distinct” em que o fluxo é associado a um único emissor. Podem no entanto ser realizadas várias reservas para emissores diferentes desde que utilizem fluxos diferentes.

SE (“shared-explicit”) – reserva da classe “shared” em que o fluxo é associado a uma lista de explicita de emissores.

WF (“wildcard-filter”) – reserva da classe “shared” que engloba todos os emissores, novos emissores que surjam são automaticamente incluídos.

As reservas da classe “distinct” não podem ser agregadas com reservas da classe “shared”. Novos estilos de reserva pode ser definidos no futuro.

Referências bibliográficas

[1] Lhamas A. (2010-2011). Aulas Teóricas de ASIST.

(...)