



Ingeniería

ISSN: 0121-750X

revista_ing@udistrital.edu.co

Universidad Distrital Francisco José de
Caldas
Colombia

Aparicio, Edith; Hernández Riaño, Velssy Liliana; Sierra Hernández, Fredy A.
Prototipo para la evaluación de calidad de servicio mediante técnicas de encolamiento en
ambientes de Backbone

Ingeniería, vol. 12, núm. 2, 2007, pp. 46-61
Universidad Distrital Francisco José de Caldas
Bogotá, Colombia

Disponible en: <http://www.redalyc.org/articulo.oa?id=498850165008>

- Cómo citar el artículo
- Número completo
- Más información del artículo
- Página de la revista en redalyc.org

redalyc.org

Sistema de Información Científica
Red de Revistas Científicas de América Latina, el Caribe, España y Portugal
Proyecto académico sin fines de lucro, desarrollado bajo la iniciativa de acceso abierto

Prototipo para la evaluación de calidad de servicio mediante técnicas de encolamiento en ambientes de Backbone

Edith Aparicio¹

Velssy Liliana
Hernández Riaño²

Fredy A. Sierra
Hernández³

Resumen

El costo en infraestructura y acceso a tecnologías emergentes en redes de telecomunicaciones, hizo necesario disponer de un planteamiento de evaluación de Calidad de Servicio utilizando mecanismos de encolamiento basado en simulación. El objetivo en este artículo fue evaluar diferentes técnicas de encolamiento modeladas bajo diversos escenarios de campus de redes *backbone*, para comparar su desempeño usando la herramienta de simulación OPNET. La evaluación apuntó hacia la medición de variables claves en el desempeño de la red como retraso, pérdida de paquetes y variación del retraso (*jitter*), entre otras. Son presentados aquí los resultados obtenidos en la utilización de éstos mecanismos al momento de otorgar diferenciación a tráfico sensible a factores de congestión presentes en las redes IP, tales como tráfico en tiempo real; obteniendo un significativo mejoramiento en el desempeño de la red.

Palabras clave: Calidad de servicio, encolamiento, OPNET, servicios diferenciados, simulación, redes troncales.

Prototype for quality of service evaluation by means of queueing techniques in backbone environments

Abstract

The cost in infrastructure and access to emerging technologies on telecommunications networks, made necessary to have a mechanism the quality of service assessment using queueing based. Objective in this article was evaluate different technical modeled various scenes of campus *backbone* network, in order to compare his performance using the tool of simulation OPNET. The evaluation aimed to the measurement of keys variables in

the performance networks like delay, packet loss and *jitter*, among others. Are show here the results once these were gotten from utilization mechanisms when granting differentiation to sensitive traffics to present factors of congestion in IP networks, such like traffic in real time; collecting a significant improvement in the network performance.

Key words: Quality of service, queuing, OPNET, differentiated services, simulation backbone.

1. Introducción

La impredecible naturaleza del tráfico creó el problema de congestión en las redes IP [1], [2], [3], [4], [7], [10]. Las aplicaciones que utilizan flujos de tiempo real o interactividad están imponiéndole a la red necesidades de tiempos de respuestas rápidos y predecibles. Los clientes que utilizan redes IP están demandando garantías de ancho de banda. Las redes *backbone* de hoy en día no cuentan con suficiente sobre aprovisionamiento de capacidad en los enlaces, debido al auge de aplicaciones hambrientas de ancho de banda [7], [8], [14]. Los enrutadores IP fueron diseñados para comprobar solamente la dirección IP del destino de una tabla de reenvío, encontrar el salto siguiente y remitir el datagrama IP. Si la cola para el salto siguiente es larga, el datagrama será retrasado. Si la cola está llena el datagrama puede incluso ser descartado. Cuando ocurre la congestión allí necesita la manera de clasificar el tráfico hacia afuera. Los paquetes que han estado marcados se pueden identificar y colocar en colas. El IETF ha desarrollado la arquitectura Diferenciación de Servicios (DiffServ) [12], [14], [20], [29], [31], [41] e integración de servicios (IntServ) [18], [19], [21], [22]-[25], [27] para administrar flujos de datos en redes IP como una solución. Sin embargo, la calidad de

¹ Directora del Grupo de Investigación GITEM, Universidad Distrital.

² Pertenencia al Grupo de Investigación SÓCRATES, Universidad de Córdoba.

³ Pertenencia al Grupo de Investigación INTERNET INTELIGENTE, Universidad Distrital.

servicio involucra diversos mecanismos: control de pérdida de paquetes cuando ocurre congestión durante un periodo de ráfaga, el establecimiento de prioridades de tráfico, dedicación de ancho de banda sobre una base por aplicación, evasión de la congestión, administración de la congestión cuando ocurre; entre otros, [5], [6], [11], [14], [33], [36], [37].

Los esquemas de Control de Congestión y Calidad de Servicio siguen siendo uno de los pilares principales para la robustez de Internet, como es mostrado por Floyd y Caída (1999a). Sin embargo, la congestión continúa siendo el obstáculo principal a la Calidad de Servicio (QoS) sobre la Internet [8]. Aunque un número de esquemas hayan sido propuestos, la investigación sobre congestión y calidad de servicio aun continúa, como lo podemos ver en los proyectos de investigación Europeos: [38] describe la evolución actual de arquitecturas de QoS, mecanismos, y protocolos en la Internet, como el curso actual del *framework* de la Unión europea consolidado en proyectos de investigación sobre redes IP premium; [39] presenta los principios y el funcionamiento básico de un esquema de control de admisión de servicio para la entrega de QoS en redes DiffServ IP; [40] define una arquitectura basada-DiffServ para entregar QoS sobre demanda a aplicaciones solicitadas; [41] define un *framework* para el aprovisionamiento de servicios de comunicación avanzados en redes IP premium caracterizadas por un alto grado de complejidad, en términos no solo de escala, sino también de número de operadores y heterogeneidad tecnológica. [11] evalúa la calidad de servicio en VoIP en redes backbone y compila los resultados de varios estudios en un solo modelo solo para evaluar Voice-over-IP (VoIP) basados en medidas de retraso y pérdidas.

En el campo de las redes de telecomunicaciones se ha experimentado un crecimiento exponencial a nivel mundial, esto ha dado lugar a la necesidad de su sofisticación. Se hizo necesario disponer de un simulador de red que ofrezca herramientas potentes con el objetivo de diseñar modelos, simular datos y analizar redes. OPNET *Modeler* es capaz de simular una gran variedad de redes; dispositivos, protocolos y aplicaciones. Proporcionando a los usuarios escalabilidad y flexibilidad, cualidades que permiten una forma efectiva de demostrar los diferentes tipos de re-

des y protocolos, para trabajar en procesos de investigación y desarrollo [16], [31].

En este artículo, nosotros damos una revisión corta del planteamiento estándar propuesto para QoS en redes IP: servicios diferenciados (DiffServ). Varios asuntos surgen al intentar llevar a cabo estas arquitecturas del mundo real en una herramienta de simulación como OPNET: modelamiento, perfiles de tráfico, configuración de servicios, recolección de estadísticas, etc. Finalmente, el artículo discute los resultados existentes y la dirección actual de investigación y desarrollo.

2. Calidad de servicio

Calidad de Servicio (QoS) admite un nivel de garantía para una aplicación tal que sus requerimientos de tráfico y servicio puedan ser satisfechos, y proveerla requiere la cooperación de elementos de la red de extremo a extremo. Cualquier garantía de QoS es solamente tan buena como el enlace más débil en la cadena entre el remitente y el receptor. No crea ancho de banda; QoS puede solamente administrar el ancho de banda de acuerdo con las demandas de aplicaciones y establecimientos de la red.

El grupo Internet Engineering Task Force (IETF) ha propuesto muchos modelos de servicio y mecanismos para cumplir con el requerimiento de QoS. Los más destacados entre ellos están el modelo de Servicios Integrados/RSVP [18], [19], el modelo de Servicios Diferenciados [20], y MPLS [45]. Nos concentramos principalmente en el modelo de Servicios Diferenciados. El objetivo fue diseñar una calidad de red consciente, que pueda soportar todo tipo de tráfico, especialmente tráfico en tiempo real [14].

2.1. Diferenciación de Servicios

Los avances de Servicios Diferenciados para el protocolo IP están dirigidos a permitir la discriminación de un servicio escalable en la Internet sin la necesidad de un estado per-flow y señalización en cada salto [22], [25]. La arquitectura de Servicios Diferenciados usa el campo de TOS (el tipo del servicio) en el encabezamiento de IPv4 y el campo Clase de Tráfico en IPv6, para especificar clases de tráfico diferentes. DiffServ minimiza la

señalización y se concentra en los flujos totales y comportamiento por salto (PHB: Per Hop Behavior) aplicable a una red con un amplio conjunto de clases de tráfico.

El campo de TOS, también llamado el DiffServ Code Point (DSCP), proporciona seis bits para el uso actual y dos bits para uso futuro. El tráfico que ingresa al dominio de la red al router de borde es clasificado, conformado y vigilado sobre la base del campo DS. Todos los paquetes que tienen el mismo requisito de campo DS son tratados como flujos agregados y consiguen el mismo tratamiento. Actualmente el Internet Engineering Task Force (IETF) ha definido tres PHB's: envío acelerado, envío seguro, y valor por defecto. La arquitectura de DiffServ usa WFQ (ponderado honradamente haciendo cola) para priorización de tráfico y Red (detección temprana aleatoria) para supervisión de tráfico.

2.2. Encolamiento

El "manejo de congestión" es un término general usado para nombrar los distintos tipos de estrategia de encolamiento que se utilizan para manejar situaciones donde la demanda de ancho de banda solicitada por las aplicaciones excede el ancho de banda total de la red, controlando la inyección de tráfico a la red, para que ciertos flujos tengan prioridad sobre otros.

- a) FIFO. Es el tipo más simple de encolamiento, se basa en el siguiente concepto: el primer paquete en entrar a la interfaz, es el primero en salir. Es adecuado para interfaces de alta velocidad, sin embargo, no para bajas, ya que FIFO es capaz de manejar cantidades limitadas de ráfagas de datos. Si llegan más paquetes cuando la cola está llena, éstos son descartados. No tiene mecanismos de diferenciación de paquetes.
- Priority Queuing. El Encolamiento de Prioridad (PQ: Priority Queueing) consiste en un conjunto de colas, clasificadas desde alta a baja prioridad. Cada paquete es asignado a una de estas colas, las cuales son servidas en estricto orden de prioridad. Las colas de mayor prioridad son siempre atendidas primero, luego la siguiente de menor prioridad y así sucesivamente. Si una cola de menor prioridad está siendo atendida, y un paquete ingresa a una cola de mayor prio-

ridad, ésta es atendida inmediatamente. Este mecanismo se ajusta a condiciones donde existe un tráfico importante, pero puede causar la total falta de atención de colas de menor prioridad.

- d) Custom Queuing. Para evadir la rigidez de PQ, se opta por utilizar Encolamiento Personalizado (CQ: Custom Queueing). Permite al administrador priorizar el tráfico sin los efectos laterales de inanición de las colas de baja prioridad, especificando el número de paquetes o bytes que deben ser atendidos para cada cola. Se pueden crear hasta 16 colas para categorizar el tráfico, donde cada cola es atendida al estilo *Round-Robin*. CQ ofrece un mecanismo más refinado de encolamiento, pero no asegura una prioridad absoluta como PQ. Se utiliza CQ para proveer a tráficos particulares de un ancho de banda garantizado en un punto de posible congestión, asegurando para este tráfico una porción fija del ancho de banda y permitiendo al resto del tráfico utilizar los recursos disponibles.
- e) Class-Based WFQ. WFQ tiene algunas limitaciones de escalamiento, ya que la implementación del algoritmo se ve afectada a medida que el tráfico por enlace aumenta; colapsa debido a la cantidad numerosa de flujos que analizar. CBWFQ fue desarrollada para evitar estas limitaciones, tomando el algoritmo de WFQ y expandiéndolo, permitiendo la creación de clases definidas por el usuario, que permiten un mayor control sobre las colas de tráfico y asignación del ancho de banda. Algunas veces es necesario garantizar una determinada tasa de transmisión para cierto tipo de tráfico, lo cual no es posible mediante WFQ, pero sí con CBWFQ. Las clases que son posibles implementar con CBWFQ pueden ser determinadas según protocolo ACL, valor DSCP, o interfaz de ingreso. Cada clase posee una cola separada, y todos los paquetes que cumplen el criterio definido para una clase en particular son asignados a dicha cola. Una vez que se establecen los criterios para las clases, es posible determinar cómo los paquetes pertenecientes a dicha clase serán manejados. Si una clase no utiliza su porción de ancho de banda, otras pueden hacerlo. Se pueden

configurar específicamente el ancho de banda y límite de paquetes máximos (o profundidad de cola) para cada clase. El peso asignado a la cola de la clase es determinado mediante el ancho de banda asignado a dicha clase.

- f) Low Latency Queuing. El Encolamiento de Baja Latencia (LLQ: Low-Latency Queueing) es una mezcla entre Priority Queueing y Class-Based Weighted-Fair Queueing. Es actualmente el método de encolamiento recomendado para Voz sobre IP (VoIP) y Telefonía IP, que también trabajará apropiadamente con tráfico de videoconferencias. LLQ consta de colas de prioridad personalizadas, basadas en clases de tráfico, en conjunto con una cola de prioridad, la cual tiene preferencia absoluta sobre las otras colas. Si existe tráfico en la cola de prioridad, ésta es atendida antes que las otras colas de prioridad personalizadas. Si la cola de prioridad no está encolando paquetes, se procede a atender las otras colas según su prioridad. Debido a este comportamiento es necesario configurar un ancho de banda límite reservado para la cola de prioridad, evitando la inanición del resto de las colas. La cola de prioridad que posee LLQ provee de un máximo retardo garantizado para los paquetes entrantes en esta cola, el cual es calculado como el tamaño del MTU dividido por la velocidad de enlace.

3. Modelamiento de la red con opnet modeler

Utilizamos diferentes modelos de topología de red para los diferentes planteamientos de QoS. El modelamiento de una red en general puede ser dividido en dos partes esenciales; modelamiento de la topología y modelamiento de tráfico. En primer lugar, será descrito el modelamiento de la topología, es decir, las interconexiones entre nodos y enlaces usados. Y se describirán los pasos para el modelamiento de tráfico para cada uno de los diferentes modelos utilizados.

3.1. Topología de la Red

Como parte de los objetivos iniciales, se decidió simular una red backbone ATM de tipo anillo, con nodos intermedios y LANs conec-

tadas en los puntos finales Figura 25. El backbone consta de Routers Juniper M5 interconectados entre sí, por medio de una conexión ATM. El Router Switch Juniper es llamado:

JN_M5_4s_a2_ge2_sl4

Este es el Juniper Router Backbone Internet M10 (JN_M5) con cuatro puertos seriales (4s), dos puertos ATM (a4), dos puertos Gigabit Ethernet (ge2) y cuatro SLIP (sl4).

En la Figura 1 se ilustra parte de la arquitectura subyacente del Router Juniper Networks M5 Internet Backbone, donde podemos notar diferentes capas reflejadas que caracterizan a este modelo nodo; la capa de aplicación sobre IP sobre la arquitectura de ATM.

Todas las figuras con excepción de la No. 3 tienen como Fuente: OPNET Technologies

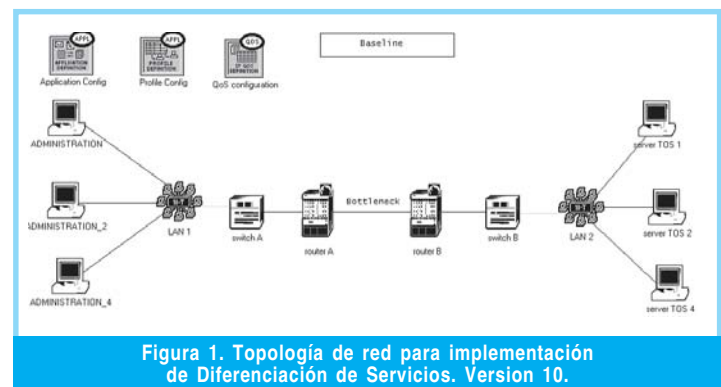


Figura 1. Topología de red para implementación de Diferenciación de Servicios. Version 10.

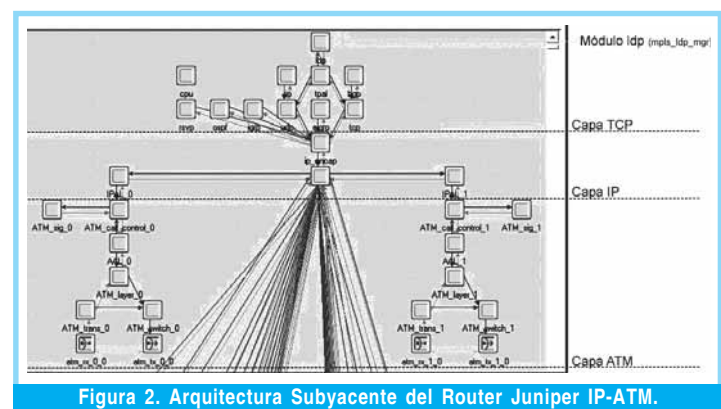


Figura 2. Arquitectura Subyacente del Router Juniper IP-ATM.

La generación de la topología de red en OPNET Modeler fue realizada manualmente, colocando los nodos individuales desde la paleta de objetos al espacio de trabajo. En primer lugar se ubicaron los routers Backbone (Juniper M10) y se unieron a través de los enlaces ATM Figura 1.

- LANs: Utilizamos la abstracción de una infraestructura de una Red de Área Lo-

cal Fast Ethernet en una topología conmutada, donde podemos modelar muchos usuarios (Ver Modelamiento del Tráfico).

- Estaciones Terminales: Empleamos el modelo nodo ethernet_wkstn_adv que representa una workstation con aplicaciones cliente-servidor ejecutándose sobre TCP/IP y UDP/IP.

Las conexiones entre los dispositivos fueron elegidas cuidadosamente para poder crear evidentes situaciones de congestión y cuellos de botella, con el fin de ver los efectos al aplicar las políticas de QoS en el sistema. Además de tener diferentes topologías de red para los diferentes esquemas de QoS, se modificó en cada experimento atributos y parámetros de los dispositivos. Una vez verificados los nodos y los enlaces de toda nuestra red, fue preciso agregarle el tráfico.

3.2. Modelamiento del Tráfico

Para modelar el tráfico de la red, debe ser modelado el comportamiento de las fuentes de tráfico, que finalmente son los usuarios. El comportamiento de un usuario en OPNET es llamado un perfil de usuario. Un perfil de usuario describe las aplicaciones usadas por el usuario, al mismo tiempo que se ocupa de especificar los patrones de tráfico seguido por las aplicaciones. De esta forma, los perfiles de usuario son especificados en los diferentes nodos de la red para generar el tráfico de la capa de aplicación.

Con el objeto Profile Configuration nosotros creamos perfiles; para tráfico de voz y para tráfico de datos, de acuerdo a los diferentes modelos de red.

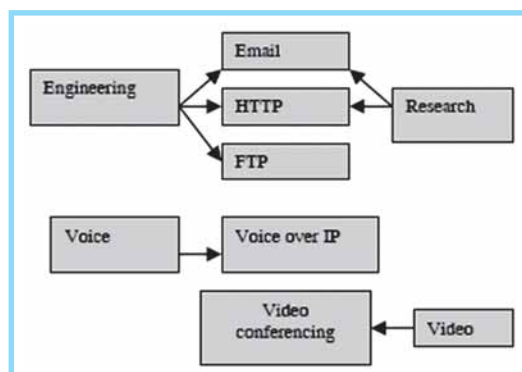


Figura 3. Ejemplo de Perfiles de usuario con sus aplicaciones. Fuente: Tomado de [46]

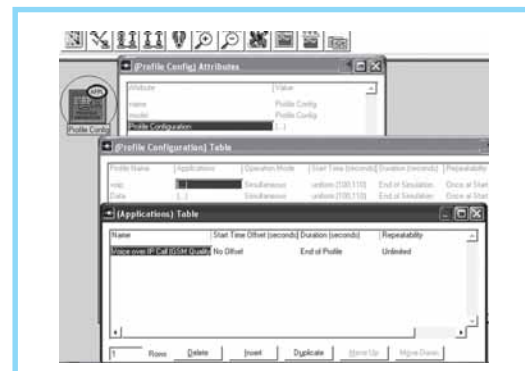


Figura 4. Definición de aplicaciones en OPNET.

Hay dos tipos de tráfico que pueden ser modelados por OPNET Modeler: tráfico explícito y el tráfico background:

Tráfico Explícito: Este tipo de tráfico es paquete por paquete, en el cual la simulación modela cada evento relacionado con el paquete (paquete creado, paquete encolado, etc.) que ocurre durante la simulación. El modelamiento del tráfico explícito proporciona los más rigurosos resultados porque modela todos los efectos del protocolo. Sin embargo, esto resulta en simulaciones más demoradas (son necesarias mas instrucciones CPU) y uso de memoria superior (porque la simulación asigna memoria para cada paquete individual). Este tipo de tráfico fue utilizado para generar el tráfico sobre las aplicaciones estándar como son FTP, Email, VoIP y Videoconferencia.

Tráfico Background: El tráfico Background afecta el desempeño del tráfico explícito al introducir retrasos adicionales. El modelo del simulador incluye los efectos del tráfico background para calcular las colas en los dispositivos intermedios y retrasos basados en la longitud de la cola, en cualquier momento durante la simulación. Sin embargo, cada paquete de este tipo de tráfico no es explícitamente modelado, por consiguiente no generará un evento en cada estado del paquete y no tiene una porción de memoria para guardar todas las características del paquete. Así pues, la simulación será más rápida y usará menos memoria.

En nuestra simulación, nosotros decidimos modelar cada computadora explícitamente en lugar de abstraer algunas partes de la topología de la red, ya que nosotros pretendemos estudiar el desempeño del backbone bajo cierto tipo de tráfico específico, como es tráfico en

tiempo real. Otro aspecto a considerar es la elección del protocolo de enrutamiento empleado por los routers del sistema. En este caso utilizamos RIP, que es el protocolo por defecto para los routers en OPNET, y consiste en elegir la ruta más corta en número de saltos.

4. Experimentos con diferenciación de servicios

El planteamiento de Servicios Diferenciados fue implementado en primera instancia sobre la red de la Figura 5, para el tráfico de Video Conferencia y Datos. De la misma manera se utilizó esta topología de red pero en vez de tráfico Video conferencia se modeló tráfico Voice over IP, cambiando las velocidades del enlace en el backbone. Finalmente, se modeló un escenario de última milla donde mostramos una topología de red y su desempeño cuando se implementan mecanismos de QoS y cuando no están habilitados estos mecanismos en la red. A continuación veremos detalles de los diferentes escenarios.

4.1 Experimento con Enlace E1 en el Backbone

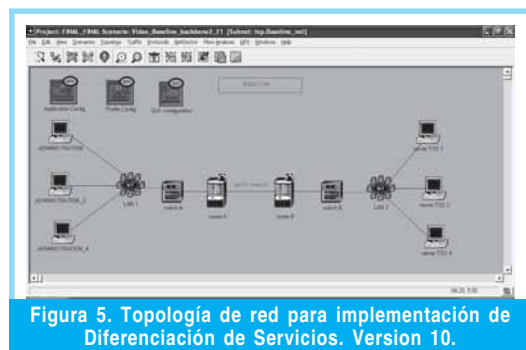


Figura 5. Topología de red para implementación de Diferenciación de Servicios. Version 10.

4.1.1 Descripción de la Topología de la red

- Dos LANs 10BaseT Ethernet
- Dos Switch Ethernet
- Dos Router Switch Juniper Router Backbone Internet M5
- Enlaces 10 BaseT Ethernet en LANs
- Enlace entre Routers Backbone ATM E1 (2.048 Mbps)

4.1.2 Descripción del Modelamiento del Tráfico

El tráfico utilizado para Diferenciación de Servicios fue Video Conferencia y se utilizó tráfico de datos para competir por los recursos.

Detalles de Aplicación. Fueron considerados tres diferentes Tipo de Servicio (ToS) para una aplicación de Videoconferencia y tres aplicaciones para datos: Email, FTP y HTTP. Los detalles de estas diferentes aplicaciones usadas se muestran en la Figura 6:

Name	Description
Email (Heavy)	(...)
File Transfer (Heavy)	(...)
Video Conferencing (streaming Multimedia)	(...)
Video Conferencing (standard)	(...)
Video Conferencing (background)	(...)
Web	(...)

Figura 6. Tabla de Definición de las Aplicaciones.

Attribute	Value
Frame Interarrival Time Information	30 frames/sec
Frame Size Information (seconds)	(...)
Symbolic Destination Name	Video Destination
Type of Service	Streaming Multimedia (4)
RSVP Parameters	None
Traffic Mix (%)	All Discrete

Figura 7. Tabla de los atributos de la conexión de Videoconferencia.

La Figura 7 muestra la descripción general de de los detalles de la conexión de Videoconferencia:

Profile Name	Applications	Operation Mode	Start Time (seconds)	Duration (seconds)	Repeatability
Background Traffic	(...)	Simultaneous	constant (100)	End of Simulation	Once at Start Time
Standard Traffic	(...)	Simultaneous	constant (100)	End of Simulation	Once at Start Time
Streaming Traffic	(...)	Simultaneous	constant (100)	End of Simulation	Once at Start Time
Campus Profile	(...)	Serial (Ordered)	constant (100)	End of Simulation	Once at Start Time

Figura 8. Configuración de perfiles usados.

Detalles de Perfil. Aquí mostramos los perfiles usados para la simulación y sus detalles. La Figura 8 muestra todos los perfiles que fueron usados:

Los cuatro perfiles que tenemos son Background Traffic, Standard Traffic, Streaming Traffic y Campus Profile. Los primeros tres son conexiones de videoconferencia con diferentes ToS. Campus Profile tiene conexiones de datos: Email, FTP y Web.

La QoS la configuramos desde el objeto global QoS Config. Desde aquí podemos esta-

blecer el esquema de encolamiento necesario para toda la red. Desde este objeto establecemos el perfil FIFO para este escenario en particular. Este perfil lo configuramos en las interfaces entre los routers donde se requiere aplicar la QoS.

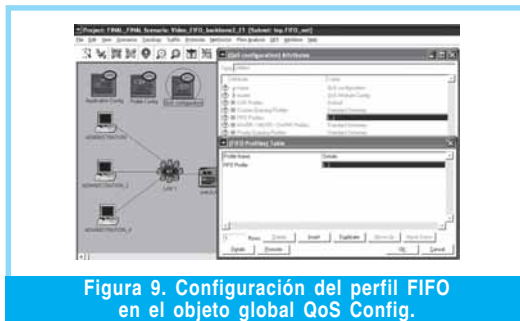


Figura 9. Configuración del perfil FIFO en el objeto global QoS Config.

Una vez configurado el perfil en el objeto global, debemos habilitar la QoS en las interfaces apropiadas de los routers. En nuestro caso configuramos la QoS de la interfaz 0 (IF0) del router A y el router B. La siguiente Figura 10 muestra la configuración de la interfaz 0 en el router A.



Figura 10. Configuración de la QoS FIFO en la interfaz IF0 del router A. Version 10.

De esta misma forma se configuraron los perfiles de encolamiento para Custom Queuing, Priority Queuing, Weighted Fair Queuing, en escenarios iguales al descrito previamente, cambiando el esquema de encolamiento. También configuramos escenarios iguales excepto en el enlace entre los routers ATM backbone, es decir comparando un enlace E1 que crea un cuello de botella, con un enlace con más capacidad como un enlace ATM DS3.

4.2 Experimento Ultima Milla

En este escenario mostramos una red usando el tipo de encolamiento WFQ y los parámetros configurados en cada nodo, usando un enlace DS0 en la última milla y un enlace DS3

en el backbone ATM. El cual comparamos con un escenario Baseline que no tiene configurado ningún tipo de esquema de QoS. A continuación mostramos las graficas del escenario Baseline y el escenario con QoS. Mostraremos en detalle la configuración del escenario con QoS, ya que el escenario Baseline no tiene mayores diferencias debido a que en OPNET las características de QoS están deshabilitadas por defecto.

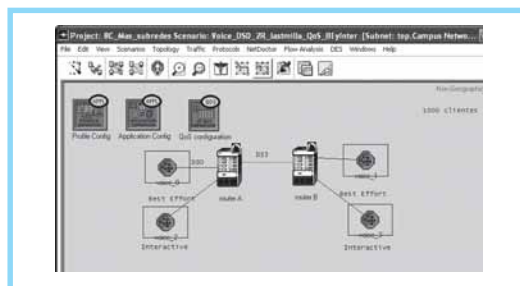


Figura 11. Topología de red con enlace DS0 en última milla y DS3 en el backbone.

En la Figura 12 observamos la topología interna de una de las subredes que componen el escenario Última milla. Cada subred consta de tres LANs que modelan 1000 usuarios de la aplicación Voice over IP.

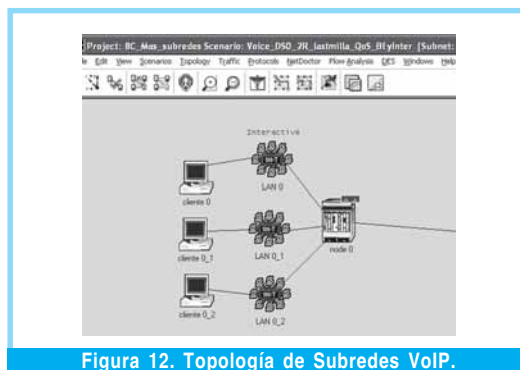


Figura 12. Topología de Subredes VoIP.

4.2.1 Descripción de la Topología de la red

- Cuatro subredes VoIP
- Dos Router Switch Juniper Router Backbone Internet M5
- Enlaces 100 BaseT Ethernet en LANs
- Enlace entre Routers Backbone ATM DS3
- Enlace última milla PPP DS0

4.2.2 Descripción del Modelamiento del Tráfico

El tráfico utilizado para este experimento fue solamente VoIP con diferente prioridad del ToS: alta y baja.

Detalles de Aplicación. Fueron considerados dos diferentes Tipo de Servicio (ToS) para una conexión de VoIP.

5. Análisis y evaluación de resultados

En esta sección mostraremos los resultados relevantes obtenidos en los diferentes experimentos simulados. Se clasifican según el tipo de aplicación y enlaces.

5.1 Video Conferencia con enlace en el backbone E1

A continuación veremos escenarios que describen un cuello de botella en el enlace backbone E1 y utilizamos tráfico de Videoconferencia.

5.1.1 Escenario Baseline: Escenario de referencia, el cual no soporta ningún tipo de encolamiento IP, y con el que se contrastaron el resto de escenarios con QoS. La Figura13 muestra las gráficas obtenidas al comparar las estadísticas de retraso extremo-extremo de paquetes (seg) y variación del retraso (seg) de los clientes ADMINISTRATION, ADMINISTRATION_2 y ADMINISTRATION_4. Se obtienen gráficas con idénticos resultados para los tres clientes.

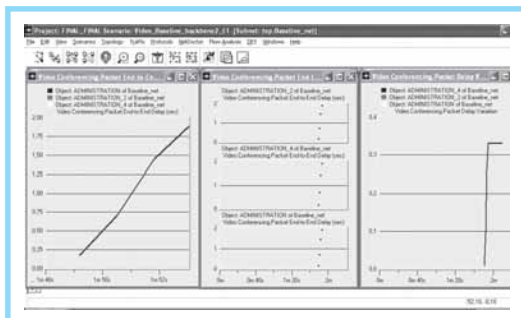


Figura 13. Retraso ETE y variación del retraso de paquetes videoconferencia con un enlace cuello de botella E1 en el backbone sin encolamiento.

5.1.2 FIFO: Escenario que soporta el perfil de encolamiento FIFO. La Figura 14 muestra las gráficas obtenidas al comparar las estadísticas de retraso extremo a extremo de paquetes (seg) de los clientes ADMINISTRATION, ADMINISTRATION_2 y ADMINISTRATION_4. Se obtienen graficas con idénticos resultados para los tres clientes. Los retrasos de retraso extremo a extremo y variación de retraso son demasiado altos. Las gráficas muestran una

tendencia creciente en el tiempo de retraso (segundos) al cabo de 2 minutos de tiempo de simulación.

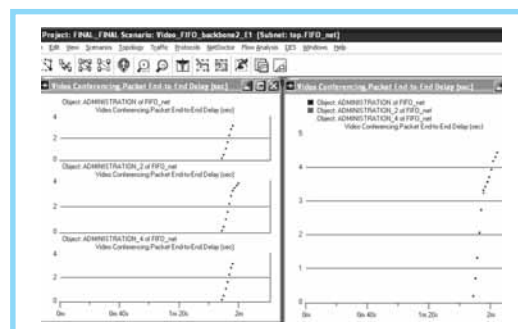


Figura 14. Retraso ETE y variación del retraso de paquetes videoconferencia con un enlace cuello de botella E1 en el backbone con encolamiento FIFO.

5.1.3 Priority Queuing (PQ): Con este tipo de encolamiento observamos resultados de estadísticas como el retraso extremo a extremo en cada una de las colas, retraso de encolamiento y tráfico descartado como se ve en la Figura15. Podemos apreciar que el tráfico descartado en la interface IP del router A (cuello de botella en el backbone) sobrepasa 1'000.000 bits/seg en la cola 0 (IF0 Q0), que fue modelada como la cola de más baja prioridad según el tipo de encolamiento Priority Queuing, mientras en las colas 1 y 3 (IF Q1 y IF Q3, mediana y alta prioridad, respectivamente) el porcentaje de tráfico descartado es inferior a 0.0 bits/seg. Los retardos extremo a extremo (ETE) (seg) de paquetes de videoconferencia empiezan con una tendencia creciente en los clientes ADMINISTRATION, ADMINISTRATION_2 y ADMINISTRATION_4, sin embargo, al cabo de 1 minuto 50 segundos de simulación, los retrasos ETE del cliente ADMINISTRATION_4, empiezan a aproximarse a 0. También podemos apreciar en estos resultados que la cola de menor prioridad tiene un mayor uso de buffer (20.000 bytes) y la variación de retraso en cola es mayor, respecto a las colas de mayor prioridad.

En la Figura 16 podemos observar que el retraso de encolamiento PQ de la cola 0 de menos prioridad, es casi cercana a 1 segundo, mientras que las colas de mayor prioridad están por debajo de 0.0.

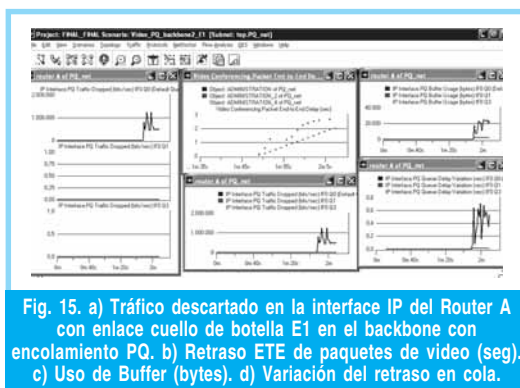


Fig. 15. a) Tráfico descartado en la interface IP del Router A con enlace cuello de botella E1 en el backbone con encolamiento PQ. b) Retraso ETE de paquetes de video (seg). c) Uso de Buffer (bytes). d) Variación del retraso en cola.

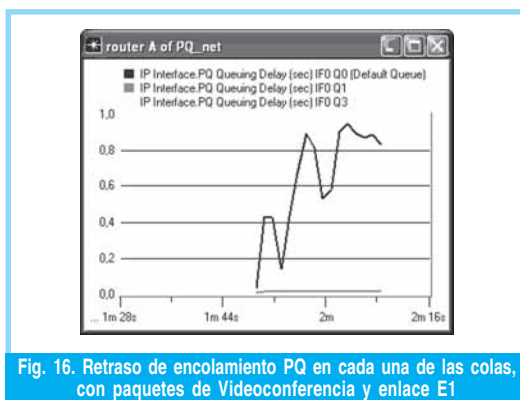


Fig. 16. Retraso de encolamiento PQ en cada una de las colas, con paquetes de Videoconferencia y enlace E1

5.1.4 Custom Queuing (CQ): Observamos los resultados en cuanto a las estadísticas Retraso extremo a extremo de paquete; tráfico enviado, uso del buffer y tráfico descartado en cada una de las colas, utilizando Custom Queuing. En la Figura 17 podemos observar que el retraso ETE en el cliente ADMINISTRATION_4, es menor que en los otros dos clientes. En cuanto al uso de buffer en la cola de menor prioridad (Q1), sobrepasa los 26.000 bytes, mientras que la cola de mayor prioridad tiene un uso de buffer 23.000 bytes. Respecto al tráfico enviado, la cola de mayor prioridad es un 50% mayor (1000'000 bits/seg), que el enviado por la cola de menor prioridad (500.000 bits/seg). El tráfico descartado en la cola de menor prioridad (800.000 bits/seg) es el triple que la de mayor prioridad (200.000 bits/seg).

5.1.5 Custom Queuing con Baja Latencia (CQ - LLQ): Configuramos en Cola de Baja Latencia la cola 0 (Q0). Podemos observar las gráficas que en los resultados de Retraso extremo a extremo de paquete, uso del buffer y

tráfico descartado la cola 0 es la que tiene menor tiempo de retraso, y el uso de buffer y tráfico descartado, son significativamente menores que las colas restantes. Figura 18.

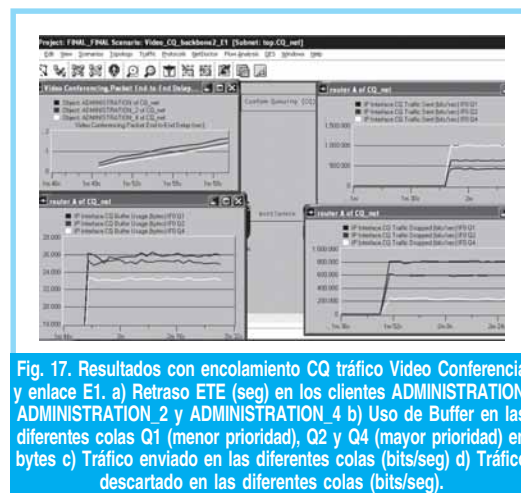


Fig. 17. Resultados con encolamiento CQ tráfico Video Conferencia y enlace E1. a) Retraso ETE (seg) en los clientes ADMINISTRATION, ADMINISTRATION_2 y ADMINISTRATION_4 b) Uso de Buffer en las diferentes colas Q1 (menor prioridad), Q2 y Q4 (mayor prioridad) en bytes c) Tráfico enviado en las diferentes colas (bits/seg) d) Tráfico descartado en las diferentes colas (bits/seg).

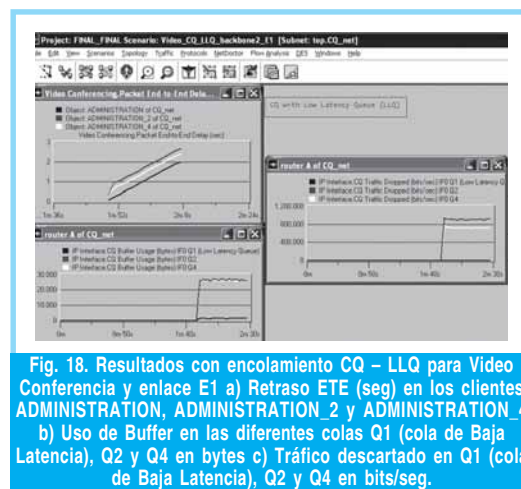


Fig. 18. Resultados con encolamiento CQ - LLQ para Video Conferencia y enlace E1 a) Retraso ETE (seg) en los clientes ADMINISTRATION, ADMINISTRATION_2 y ADMINISTRATION_4 b) Uso de Buffer en las diferentes colas Q1 (cola de Baja Latencia), Q2 y Q4 en bytes c) Tráfico enviado en Q1 (cola de Baja Latencia), Q2 y Q4 en bits/seg. d) Tráfico descartado en las diferentes colas (bits/seg).

5.1.6 Weighted Fair Queuing (WFQ): Observamos las estadísticas retraso ETE de paquetes, retraso de encolamiento en las colas Q1, Q2 y Q4 en segundos; tráfico enviado en cada una de las diferentes colas en paquetes/seg. Además un comparativo entre PQ y WFQ en cuanto a tráfico descartado. Es evidente la mejora en el tráfico descartado cuando tenemos encolamiento WFQ.

5.1.7 WFQ con Baja Latencia: Resultados de las estadísticas retraso ETE de paquetes en cada uno de los clientes y retraso de encolamiento en las colas Q1, Q2 y Q4.

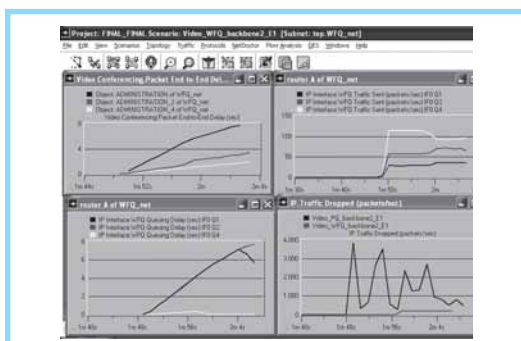


Fig. 19. Resultado WFQ Video Conferencia E1. a) Retraso ETE (seg) en los clientes ADMINISTRATION, ADMINISTRATION_2 y ADMINISTRATION_4 b) Retraso de encolamiento en el router A (cuello de botella) en las colas Q1, Q2 y Q4 en seg. c) Tráfico enviado en las diferentes colas Q1, Q2 y Q4 en paquetes por segundo d) Comparación entre el tráfico descartado en el escenario con encolamiento PQ y WFQ en paquetes/seg.

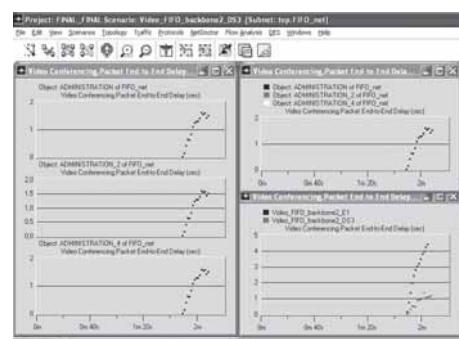


Figura 21. Resultados con encolamiento FIFO para tráfico Video Conferencia y enlace DS3 en el backbone a) Retraso ETE (seg) en los clientes ADMINISTRATION, ADMINISTRATION_2 y ADMINISTRATION_4 b) Comparación del retraso ETE entre escenarios con encolamiento FIFO pero con enlaces en el backbone E1 y DS3 en seg.

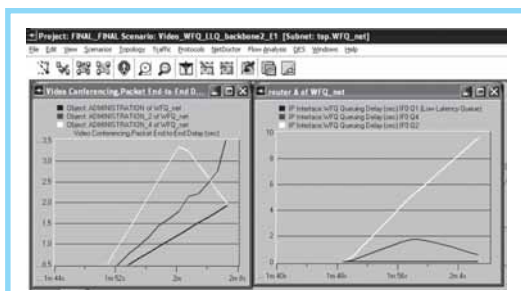


Fig. 20. Resultados con encolamiento WFQ LLQ para Video Conferencia con enlace E1. a) Retraso ETE (seg) en los clientes ADMINISTRATION, ADMINISTRATION_2 y ADMINISTRATION_4 b) Retraso de encolamiento en Q1 (cola de baja latencia), Q2 y Q4.

5.2 Tráfico de Video Conferencia con enlace DS3 en el backbone

En estos escenarios el objetivo es comparar el desempeño de la red cuando tenemos enlaces en el backbone con mayor capacidad (DS3), y aplicando los diferentes tipos de encolamiento.

5.2.1 FIFO: Obtenemos el retraso extremo a extremo de paquetes (seg). Como podemos apreciar en la Figura 20 el retraso ETE es idéntico entre los tres clientes. Hacemos un comparativo de los escenarios con el mismo tipo de encolamiento FIFO, pero variando el enlace en el backbone, es decir, un escenario con enlace cuello de botella E1, y el otro con enlace DS3 en el backbone. Podemos notar que al cabo de dos minutos de tiempo de simulación, el escenario con enlace E1 consigue un retraso de casi 5 segundos, mientras el escenario con enlace DS3, consigue un retraso inferior a 1 segundo y su tendencia es decreciente.

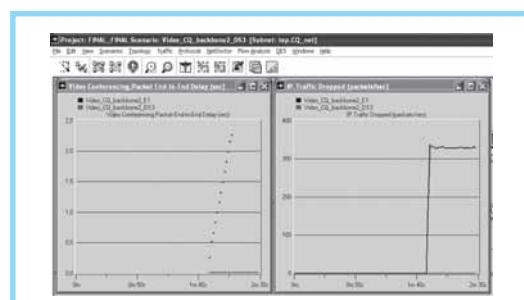


Figura 22. Comparativo con encolamiento CQ con tráfico Video Conferencia y enlaces E1 y DS3 a) Retraso ETE en segundos b) Tráfico descartado en paquetes/seg.

5.2.3 Custom Queuing con Baja Latencia (CQ-LLQ) con enlace DS3 en el backbone: En la Figura 23 se muestran los resultados de la comparación entre escenarios con enlaces E1 y DS3 de las estadísticas retraso extremo a extremo de paquetes (seg), variación del retraso de paquete y tráfico descartado.

5.2.4 Encolamiento WFQ y enlace DS3: En la Figura 24, observamos los resultados de las estadísticas retraso extremo a extremo de pa-

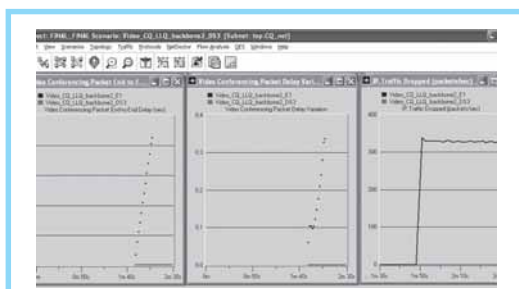


Figura 23. Resultados con encolamiento CQ – LLQ para Video Conferencia y enlaces E1 y DS3. a) Comparación del retraso ETE entre escenarios con encolamiento.

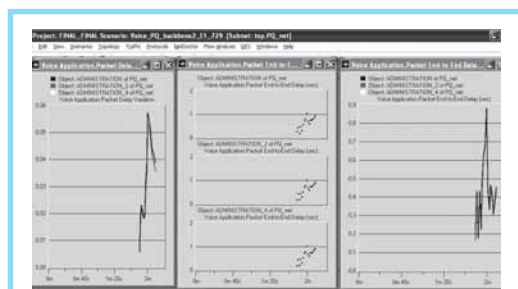


Figura 27. Resultados con encolamiento PQ enlace ATM E1 y tráfico Voice over IP en los clientes a) Variación del retraso de paquetes b) Retraso ETE de paquetes. Fuente: OPNET Technologies

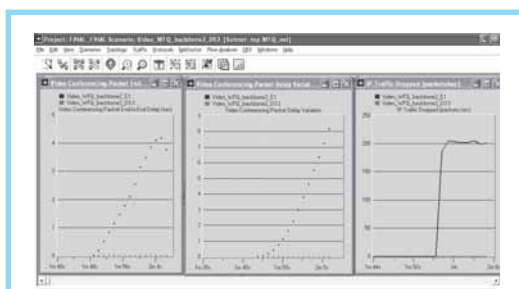


Fig. 24. Resultados comparativos con encolamiento WFQ para Video Conferencia y enlaces E1 y DS3. a) Comparación del retraso ETE entre escenarios con encolamiento WFQ pero con enlaces en el backbone E1 y DS3 en seg b) Variación del retraso c) Tráfico Descartado.

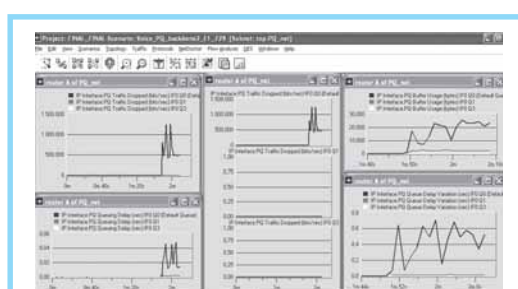


Fig. 28. Resultados en el router A con encolamiento PQ enlace ATM E1 y tráfico Voice over IP a) tráfico descartado en las diferentes colas b) Retraso de encolamiento d) Uso del buffer e) Variación del retraso en las diferentes colas.

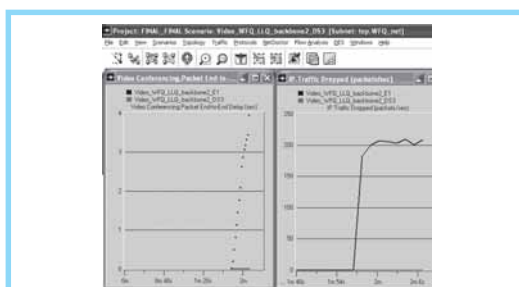


Figura 25. Resultados con encolamiento WFQ – LLQ con Video Conferencia y enlaces E1 y DS3 a) Retraso ETE b) Tráfico descartado en la interface IP.

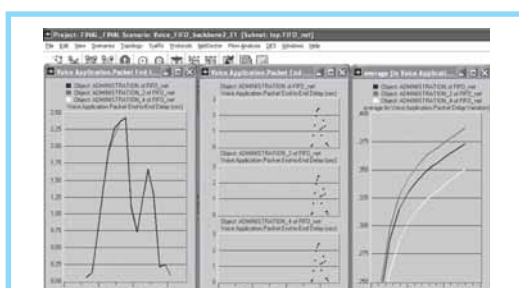


Figura 26. Resultados con encolamiento FIFO con enlace ATM E1 con tráfico Voice over IP. a) Retraso ETE en los clientes ADMINISTRATION, ADMINISTRATION_2 y ADMINISTRATION_4 b) Variación del retraso de paquetes en los diferentes clientes.

quete (seg), variación del retraso de paquete y tráfico descartado, al comparar escenarios con enlaces E1 y DS3.

5.2.5 Encolamiento WFQ con Baja Latencia DS3: Comparación entre E1 y DS3 de las estadísticas retraso ETE de paquete (seg), variación del retraso de paquete y tráfico descartado. Figura 25.

5.3 Voice over IP con enlace en el backbone E1

En los siguientes escenarios se utiliza tráfico VoIP y un enlace cuello de botella E1 en el backbone. Se compara el desempeño de la red aplicamos los diferentes tipos de encolamiento.

5.3.1 FIFO: En la Figura 26 se muestra el retraso extremo a extremo de paquetes (seg) y la variación del retraso (seg) para la aplicación voice over IP con un enlace ATM E1 en el backbone.

5.3.2 Priority Queuing (PQ): En este escenario observamos las estadísticas variación del retraso extremo a extremo (seg) y retraso extremo a extremo (seg) para la aplicación *Voice over IP* en la Figura 27 y en la Figura 28 obser-

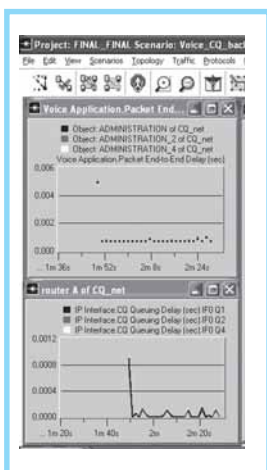


Figura 29. Resultados con encolamiento CQ enlace ATM E1 y tráfico Voice over IP a) Retraso ETE en los clientes b) Variación del retraso en las colas.

vamos estadísticas a nivel del nodo *router A*: tráfico descartado, retraso de encolamiento.

5.3.3 Custom Queuing (CQ): La Figura 29 muestra los resultados para retraso extremo a extremo de paquetes (seg), y estadísticas en el nodo router A, a nivel de la Interface IP: uso del buffer (bytes), tráfico enviado (bytes/seg) y tráfico descartado (bytes/seg) para cada cola.

5.3.4 Custom Queuing con Baja Latencia (CQ - LLQ): Este escenario es configurado 5.4.1 FIFO DS3: Compara el actual escenario DS3 el escenario Voice over IP con enlace E1, al utilizar el esquema de encolamiento FIFO, para las estadísticas retraso extremo a extremo de paquetes (seg) y uso del buffer (bytes) y tráfico enviado en cada cola (bytes/seg). Figura 32

5.3.5 Weighted Fair Queuing (WFQ): Este escenario nos muestra el retraso extremo a extremo de paquetes (seg) para la aplicación Voice over IP a nivel global. A nivel de nodos tenemos las estadísticas para el router A en la interface IP: variación del retraso de encolamiento (seg). Figura 31.

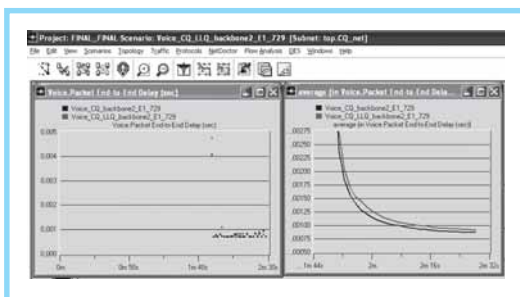


Figura 30. Retraso ETE con encolamiento CQ-LLQ enlace ATM E1 y tráfico Voice over IP.

ATM E1 y Voice over IP. a) Retraso ETE en los clientes b)Variación del retraso en cada cola. OPNET Technologies

5.4 Voice over IP con enlace en el backbone ATM DS3

En estos escenarios el objetivo es comparar el desempeño de la red cuando tenemos enlaces en el backbone con mayor capacidad (DS3) y el tipo de tráfico continúa siendo Voice sobre IP, y aplicando los diferentes tipos de encolamiento.

5.4.1 FIFO DS3: Compara el actual escenario DS3 el escenario Voice over IP con enlace E1, al utilizar el esquema de encolamiento FIFO, para las estadísticas retraso extremo a extremo de paquetes (Seg) y uso del buffer (bytes) y tráfico enviado en cada cola (bytes/seg). Figura 32

5.4.2 PQ enlace DS3: Compara este mismo escenario pero con enlace backbone E1 Voice over IP con las estadísticas a nivel de nodo router A: variación del retraso de cola (seg) para la cola Q1 y Q3, retraso de encolamiento (seg). Figura 33

5.4.3 Encolamiento CQ y enlace DS3: en la Figura 33 se compara este escenario con otro igual pero con la diferencia del enlace en el backbone E1 para la estadística retraso de paquetes extremo a extremo (seg) y variación de retraso de paquetes (seg).

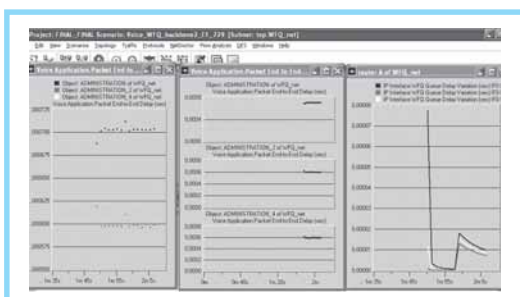


Figura 31. Resultado con encolamiento WFQ enlace ATM E1 y Voice over IP. a) Retraso ETE en los clientes b)Variación del retraso en cada cola.

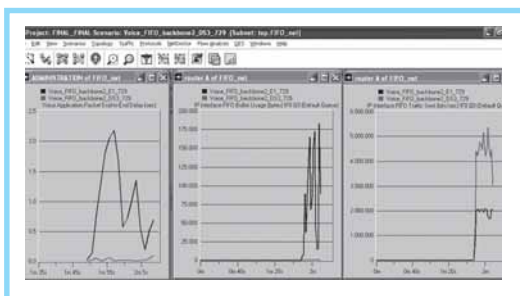


Figura 32. Resultado comparación del encolamiento FIFO para VoIP y enlaces E1 y DS3 a) Retraso ETE b) Uso de Buffer c) Tráfico Enviado.

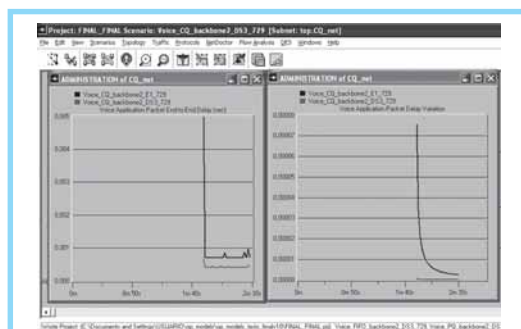


Figura 33. Resultado de comparar el encolamiento CQ para tráfico VoIP y enlaces E1 y DS3 a) Retraso ETE b) Variación del retraso.

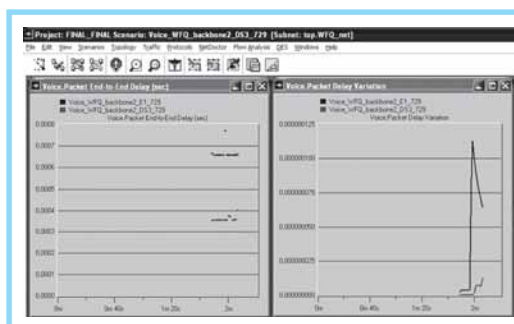


Figura 34. Resultado de comparar el encolamiento WFQ VoIP y enlaces E1 y DS3 a) Retraso ETE b) Variación del retraso. Fuente: OPNET Technologies

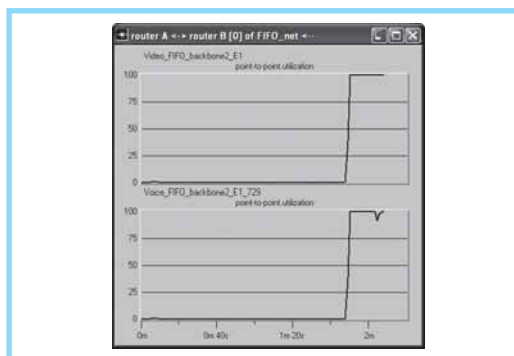


Figura 35. Resultado de la utilización del enlace con encolamiento FIFO, enlace E1 y tráfico de Voice over IP y video.

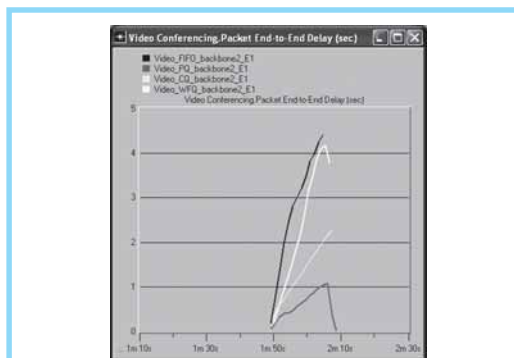


Figura 36. Resultado de retraso ETE al comparar los diferentes esquemas de encolamiento usados con tráfico de video enlace

5.4.4 Encolamiento WFQ y enlace DS3: Muestra la comparación para un enlace E1 y enlace DS3 en el backbone para las estadísticas retraso de paquete extremo a extremo (seg) y variación de retraso de paquete (seg). Figura 35

5.4.5 Utilización del enlace punto a punto: Voice over IP y Video Conferencia

En la Figura 36, observamos la utilización del enlace al utilizar el mismo esquema de encolamiento (FIFO) y variando el tipo de tráfico (Voice over IP y Video)

En la Figura 37, comparamos los diferentes tipos de encolamiento en escenarios idénticos con enlaces cuello de botella E1. Se observa que la el tiempo de retraso ETE es 0 para el tipo de encolamiento PQ y, los más altos retrasos están en el escenario con encolamiento FIFO.

En la Figura 38, comparamos observamos que la variación de retraso de paquetes cuando se utiliza el encolamiento WFQ es la mayor, mientras que la menor es la obtenida al utilizar el tipo de encolamiento CQ.

En la Figura 39, comparamos los diferentes tipos de encolamiento y obtuvimos la estadística tráfico descartado, donde se puede apreciar que el mayor porcentaje de tráfico descartado se da al utilizar los tipos de encolamiento PQ y FIFO (4000 y 3000 paquetes por segundo), mientras que los mejores resultados se obtienen al utilizar el tipo de encolamiento WFQ.

5.5 Escenario Última Milla (VoIP)

En la Figura 40 comparamos los dos tipos de QoS (Baja Prioridad y Alta prioridad) en dos clientes para la estadística retraso ETE de paquetes VoIP. El cliente con baja prioridad (voice_0) tiene retrasos mayores a 3 segundos, mientras el cliente de alta prioridad conserva retrasos inferiores a 0.0 segundos. De Igual forma notamos que el tráfico enviado en la interface IP para la cola de mayor ToS es superior a 125 paquetes/seg, mientras que en la cola de más baja prioridad, el número de paquetes enviados es de aproximadamente 75 paquetes/Seg.

En la Figura 41 observamos 150.000 bits/seg de utilización del enlace de la red frente a la mínima utilización cuando se utiliza ToS con baja prioridad.

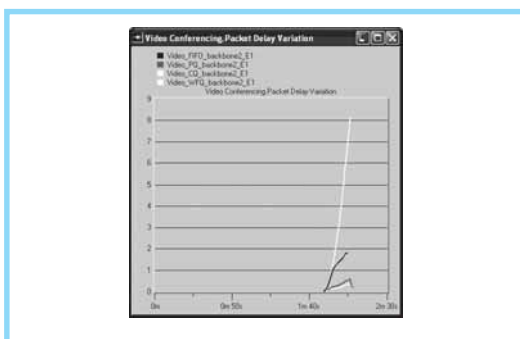


Figura 37. Resultado de variación del retraso al comparar los diferentes esquemas de encolamiento usados con tráfico de video enlace E1.

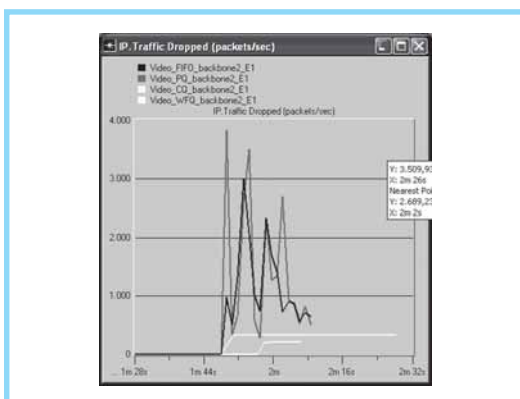


Figura 38. Resultado de tráfico descartado al comparar los diferentes esquemas de encolamiento usados con tráfico de video enlace E1.

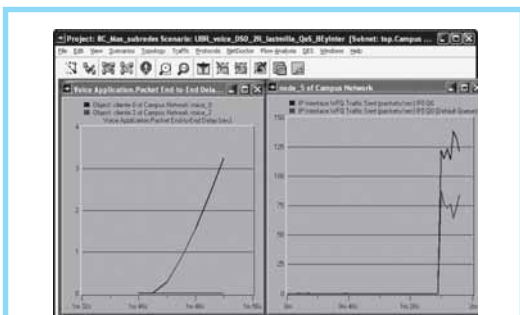


Figura 39. Resultados del escenario Ultima milla con ToS alto y bajo (Best effort e Interactive) con tráfico VoIP a) Retraso ETE (seg) b) Tráfico enviado en las dos colas con ToS 0 y ToS 6.

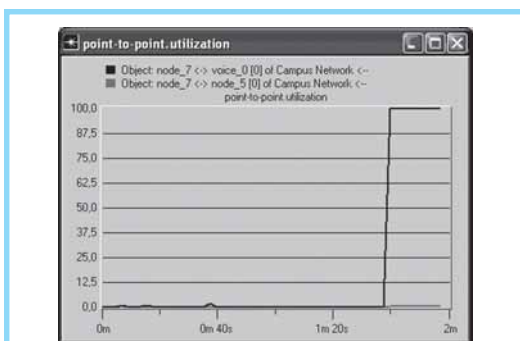


Figura 40. Resultado de Utilización (bits/seg) con tráfico VoIP escenario Ultima milla.

En la Figura 42 observamos el retraso de encolamiento para los dos tipos de QoS (Baja Prioridad y Alta prioridad).

En la Figura 43 observamos el throughput en bits/Seg para los dos tipos de QoS (Baja Prioridad y Alta prioridad).

6. Conclusiones

OPNET es un poderoso software de simulación de redes bastante usado en la investigación en redes de Telecomunicaciones. Un problema experimentado con este software de simulación es su complejidad, es decir, es necesario especificar un gran conjunto de parámetros para correr una simulación; debido a que es una herramienta desarrollada principalmente para recrear redes reales, implementando escenarios reales tipo “what-if” para encontrar qué pasa con el desempeño de una red en caso de cambiar el diseño de una red, se hace necesario incrementar el número de usuarios o cambiar patrones de tráfico.

Inicialmente se estudió el algoritmo FIFO (primero en entrar, primero en salir) que es el algoritmo utilizado por defecto en la transmisión. Debido a que no toma decisiones sobre la prioridad de los paquetes, el tráfico multimedia va a ser tratado de la misma forma que el tráfico de datos tradicionales, lo que no proporciona ningún nivel de calidad. En consecuencia será necesario utilizar otros algoritmos más sofisticados para las redes inteligentes actuales.

Así, cuando necesitemos dar el mayor nivel de prioridad al tráfico que consideremos más importante según el protocolo de red utilizado, el algoritmo más adecuado será PQ (Priority Queuing) que dispone de diferentes colas teniendo asignada cada una de ellas una prioridad de mayor a menor.

CQ (Custom Queuing) por su parte tiene la ventaja de que permite garantizar la asignación de un ancho de banda en aquellos puntos donde se produce la congestión, dejando el resto para cualquier otro tipo de tráfico. Sin embargo, y pese a la mejora, estos dos últimos algoritmos utilizan una configuración estática y no se adaptan a los requisitos cambiantes de las redes, de ahí la utilización de WFQ

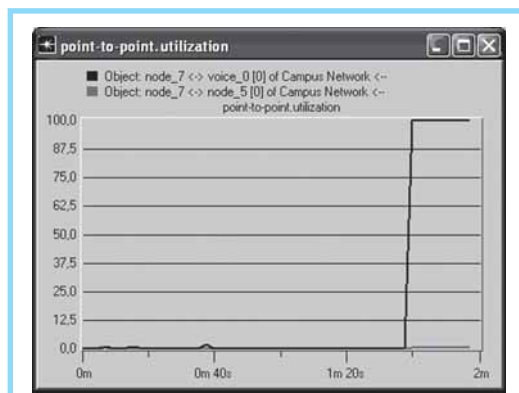


Figura 41. Resultado VoIP Última milla Retraso de encolamiento.

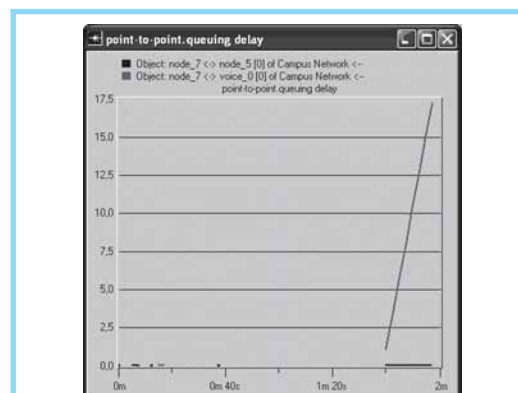


Figura 42. Resultado VoIP Última milla Throughput.
Fuente: OPNET Technologies

(Weighted Fair Queuing). Este es el más adecuado para situaciones en las que se necesita proporcionar un tiempo de respuesta consistente a cualquier tipo de usuarios sin necesidad de aumentar de forma excesiva el ancho de banda, asegurándose siempre la existencia del mismo, proporcionando así un servicio predecible. De esta manera, el tráfico interactivo será situado al principio de la cola, reduciendo así el tiempo de respuesta y pudiendo compartir después el resto de ancho de banda entre flujos de tráfico de baja prioridad.

Dentro de los protocolos de calidad de servicio existentes nos encontramos con el protocolo de reserva de recursos o RSVP que permite a las aplicaciones solicitar una calidad de servicio específica a la red lo que va a garantizar un nivel de QoS para los mismos. Esto le convierte actualmente en el protocolo que proporciona un mayor nivel de calidad de servicio, implementándose el control tanto en las aplicaciones como en la misma red. Sin embargo, para conseguirlo, es necesario utilizar técnicas de señalización que conlleven decisiones de encaminamiento y que cargarán aún más la red, término a tener en cuenta, sobre todo cuando este protocolo se aplique sobre IP. Además, la dificultad en la escalabilidad y el hecho de no poder trabajar sobre redes no RSVP, han provocado la aparición de otras alternativas.

En primer lugar, cabe destacar el modo de transferencia asíncrono. ATM es una tecnología que soporta con altos niveles de garantía la tendencia actual de generación de tráfico

multimedia, pues ha sido concebida desde su nacimiento como una tecnología capaz de proporcionar altos niveles de calidad de servicio.

Referencias bibliográficas

- [1] Zefeng Nia, Zhenzhong Chena, King Ngi Ngana, *A real-time video transport system for the best-effort Internet*. Signal Processing: Image Communication 20 (2005) pag. 277–29.
- [2] Aurelio La Corte, Sabrina Sicari, *Assessed quality of service and voice and data integration: A case study*. Computer Communications 29 (2006) pag. 1992–2003.
- [3] Spyridon L. Tompros, Spyridon Denazis, *Interworking of heterogeneous access networks and QoS provisioning via IP multimedia core networks*. Computer Networks 52 (2008) pag. 215–227.
- [4] Hans van den Berg, Michel Mandjes, Remco van de Meent, *QoS-aware bandwidth provisioning for IP network links*. Computer Networks 50 (2006) 631–647.
- [5] J.B. Pippas, I.S. Venieris, *Applying delay random early detection to IP gateways*. Computer Communications 24 (2001) pag. 1370–1379.
- [6] Christos Bouras, Afrodite Sevasti, *An analytical QoS service model for delay-based differentiation*. Computer Networks 51 (2007) pag. 3549–3563.
- [7] Sheau-Ru Tong, Chun-Cheng Chang, *Harmonic DiffServ: Scalable support of IP multicast with Qos heterogeneity in DiffServ backbone networks*. Computer Communications 29 (2006) pag. 1780–1797.
- [8] Athina Markopouloua, Fouad Tobagib, Mansour Karamb, *Loss and Delay Measurements of Internet Backbones*. Computer Communications 29 (2006) pag. 1590–1604.
- [9] A. Bak, W. Burakowski, F. Ricciato, S. Salsano, H. Tarasiuk, *A framework for providing differentiated QoS guarantees in IP-based network*. Computer Communications 26 (2003) pag. 327–337.
- [10] S. Georgoulas, P. Trimintzios, G. Pavlou, K. Ho, *An integrated bandwidth allocation and admission control framework for the support of heterogeneous real-time traffic in class-based IP networks*. Computer Communications 31 (2008) pag. 129–152.
- [11] Vitalio Alfonso Reguera, Félix F. Álvarez Paliza, Walter Godoy Jr., Evelio M. García Fernández, *On the impact of active queue management on VoIP quality of service*. Computer Communications 31 (2008) pag 73–87.

- [12] Eun-Hee Cho, Kang-Sik Shin, Sang-Jo Yoo, *SIP-based QoS support architecture and session management in a combined IntServ and DiffServ networks*. Computer Communications 29 (2006) pag. 2996–3009.
- [13] Irfan Awan, Shakeel Ahmad, Bashir Ahmad, *Performance analysis of multimedia based web traffic with QoS constraints*. Journal of Computer and System Sciences 74 (2008) pag. 232–242.
- [14] Sergio Herreria-Alonso, Andrés Suarez-González, Manuel Fernández-Veiga, Raúl F. Rodríguez-Rubio, Cándido López-García, *Improving aggregate flow control in differentiated services networks*. Computer Networks 44 (2004) pag. 499–512.
- [15] Irfan Awan, Shakeel Ahmad, Bashir Ahmad, *Performance analysis of multimedia based web traffic with QoS constraints*. Journal of Computer and System Sciences 74 (2008) pag. 232–242.
- [16] K. Salah, P. Callyam, M.I. Buhari, *Assessing readiness of IP networks to support desktop videoconferencing using OPNET*. Journal of Network and Computer Applications, Article in Press.
- [17] The QoS Forum. <http://www.qosforum.com>.
- [18] W. Almesberger. *Scalable Resource Reservation for the Internet*. PhD thesis, EPFL, Nov 1999.
- [19] T. Ferrari. *QoS Support for Integrated Networks*. PhD thesis, University of Bologna, Nov 1998.
- [20] S. Blake, D. Black, M. Carlson, E. Davies, Z. Wang, and W. Weiss. *An Architecture for Differentiated Services*. RFC 2475, IETF Network Working Group, December 1998.
- [21] R. Braden, D. Clark, and S. Shenker. *Integrated Services in the Internet Architecture: an Overview*. RFC 1633, ISI, MIT and PARC, June 1994.
- [22] IETF. Integrated Services (Intserv) working group. <http://www.ietf.org/html.charters/intserv-charter.html>
- [23] J. Wroclawski, *Specification of controlled-load network element service*. RFC 2211, Sept. 1997
- [24] S. Shenker, C. Partridge, and R. Guerin, *Specification of guaranteed quality of service*. RFC 2212, Sept 1997.
- [25] R. Braden, L. Zhang, S. Berson, S. Herzog, and S. Jamin, *Resource reservation protocol (RSVP)-Version 1, functional specification*. RFC 2205, Sept 1997
- [26] Garcia, L. *Communications Networks*, Ed. McGraw Hill, 2000.
- [27] Gene Gaines, Marco Festa, *A survey of RSVP/QoS Implementations*. RSVP Working Group, July 1998.
- [28] www.qosforum
- [29] K. Nichols, S. Blake, F. Baker, and D. Black. *Definition of the Differentiated Services Field (DS Field) in the IPv4 and IPv6 Headers*. RFC 2474, IETF Network Working Group, December 1998
- [30] Kucheria, Amit P. *Scalable Emulation of IP Networks through Virtualization*, Tesis de Maestría, Universidad de Kansas, 2003
- [31] Jun Wang, Klara Nahrstedt, Yuxin Zhou, *Design and Implement Differentiated Service Routers in OPNET*. University of Illinois at Urbana-Champaign Department of Computer Science.
- [32] C. Aras, J. Kurose, D. Reeves, and H. Schulzrinne, *Real-time communication in packet switched networks*. Proceedings of the IEEE, vol. 82, pp. 122-139, January 1994.
- [33] S. Floyd and V. Jacobson, *Random Early Detection Gateways for Congestion Avoidance*. IEEE/ACM Transactions on Networking, 1(4), August 1993.
- [34] The ATM Forum, *Traffic Management Specification*, Version 4.0, February 1996.
- [35] Stallings, William, *Redes e Internet de Alta Velocidad: Rendimiento y Calidad de Servicio*. 2ª Ed. Madrid, Prentice-Hall, 2004
- [36] T. Ferrari: *QoS support for Integrated Networks*. PhD thesis, University of Bologna, Nov 1998.
- [37] Omer Dedecoglu, *Impacts of RED in a Heterogeneous Environment*. Master thesis, Albuquerque, New Mexico May, 2003.
- [38] S. Giordano et al., *Advanced QoS Provisioning in IP Networks: The European Premium IP Projects*. IEEE Commum. Mag., vol. 41, no. 1, January 2003
- [39] E. Mykoniati et al., *Admission Control for Providing QoS in Diffserv IP Networks: The Tequila Approach*. IEEE Commum. Mag., vol. 41, no. 1, January 2003
- [40] T. Engel et al., *AQUILA: Adaptive Resource Control for QoS Using an IP-Based Layered Architecture*. IEEE Commum. Mag., vol. 41, no. 1, January 2003
- [41] IETF Internet Draft, M. Pullen, *Expanded Simulation Models for IntServ and DiffServ with MPLS*. November, 2002
- [42] M. Potts, *QoS: Quality of Service for IP networks*. Del IST-2000-26418, NGN Initiative, February 2002
- [43] G. Cortese et al., *CADENUS: Creation and Deployment of End-User Services in Premium IP Networks*. IEEE Commum. Mag., vol. 41, no. 1, January 2003
- [44] G. Huston, *Next Steps for the IP QoS Architecture*. RFC 2990, November 2000
- [45] Desire Oulai, Steven Chamberland, Samuel Pierre, *Routing and admission control with multiconstrained end-to-end quality of service in MPLS networks*. Computer Communication (2008)
- [46] Jason Schreiber, Mehrdad Khodai Joopari, M. A. Rashid, *Performance of voice and video conferencing over ATM and Gigabit Ethernet backbone networks*. Res. Lett. Inf. Math. Sci. (2005) vol. 7, pag. 19-27

Edith Aparicio

Ph.D. En Ciencias Técnicas Universidad Central de las Villas UCLV. Maestría en Teleinformática Universidad Distrital Francisco José de Caldas. Especialización en Gerencia de Proyectos Educativos en la misma universidad. Perteneció al grupo de investigación GITEM. Actualmente es Directora de la Maestría en Ciencias de la Información y la Comunicación. medicina@udistrital.edu.co

Velssy Liliana Hernández

Obtuvo su título de Maestría en Teleinformática en la Universidad Distrital Francisco José de Caldas, Ingeniera de Sistemas de la misma universidad. Actualmente es docente de la Universidad de Córdoba Facultad de Ingeniería de Sistemas. Perteneció al grupo de investigación SÓCRATES. velssyliliana@gmail.com

Fredy A. Sierra

Maestría en Teleinformática de la Universidad Distrital Francisco José de Caldas (c). Obtuvo su título en Ingeniería Electrónica en la misma universidad. Perteneció al grupo de investigación Internet Inteligente. Actualmente es instructor del SENA. intecman@gmail.com