

Universidad Central “Marta Abreu” de Las Villas

Facultad de Ingeniería Eléctrica

Departamento de Electrónica y Telecomunicaciones



TRABAJO DE DIPLOMA

**Propuesta de un backbone MPLS para la red NGN en Villa
Clara.**

Autor: Reinier Tirado de la Cruz

Tutor: Ing. Rafael Olivera Solís

Santa Clara, Cuba

2015

"Año 57 de la Revolución"

Universidad Central “Marta Abreu” de Las Villas

Facultad de Ingeniería Eléctrica

Departamento de Electrónica y Telecomunicaciones



TRABAJO DE DIPLOMA

Propuesta de una red NGN para la provincia de Villa Clara.

Autor: Reinier Tirado de la Cruz

rtirado@uclv.edu.cu

Tutor: Ing. Rafael Olivera Solís

rolivera@uclvedu.cu

Santa Clara, Cuba

2015

"Año 57 de la Revolución"



Hago constar que el presente trabajo de diploma fue realizado en la Universidad Central “Marta Abreu” de Las Villas como parte de la culminación de estudios de la especialidad de Ingeniería en Telecomunicaciones y Electrónica, autorizando a que el mismo sea utilizado por la Institución, para los fines que estime conveniente, tanto de forma parcial como total y que además no podrá ser presentado en eventos, ni publicados sin autorización de la Universidad.

Firma del Autor

Los abajo firmantes certificamos que el presente trabajo ha sido realizado según acuerdo de la dirección de nuestro centro y el mismo cumple con los requisitos que debe tener un trabajo de esta envergadura referido a la temática señalada.

Firma del Tutor

Firma del Jefe de Departamento
donde se defiende el trabajo

Firma del Responsable de
Información Científico-Técnica

PENSAMIENTO

“Hay una fuerza motriz más poderosa que el vapor, la electricidad y la energía atómica: la voluntad”.

Albert Einstein

DEDICATORIA

Quiero dedicarle este trabajo a mi familia en especial:

*A mi mamá por apoyarme durante mi largo tiempo de estudio en todo lo que me
hizo falta;*

A mi hermana;

*A mi novia Rosmely por estar a mi lado incondicionalmente y ayudarme
siempre;*

*A todos mis amigos en especial a esa manada que tanto quiero formada por
Ernesto, el Wao, la Bala, el Robe, Pedrito, el Chonguito, el Gavi, Yoanner y
Leandro destacando de manera muy especial a Ernesto y a su novia Zulmary
por dedicar tantas horas de su tiempo para lograr este trabajo;*

A mis amigos Erik y su esposa Gretter;

A Eduardito y su Familia.

AGRADECIMIENTOS

A mi mamá por estar siempre cuando la necesité durante mi época de estudiante a mi hermana.

A todas las personas que de una forma u otra influyeron en mi formación profesional como ingeniero en telecomunicaciones y electrónica.

A todos mis compañeros de aula por compartir todos estos años juntos y regalarme la mejor etapa de mi vida.

A mi novia por apoyarme siempre, aun cuando la tarea era difícil.

A mi tío por ser un ejemplo a seguir durante toda mi vida.

A toda la manada de tele.

A mi amigo Ernesto por dedicarle su tiempo a este trabajo, sin él, nada de esto sería posible y también a su novia por su apoyo y su tiempo.

A mi tutor por guiarme y ayudarme en este trabajo.

A Jose por su apoyo en la impresión de la tesis logrando de esta forma que me quede un recuerdo importante para toda mi vida.

TAREA TÉCNICA

1. Búsqueda de información sobre las Redes de Próxima Generación.
2. Descripción de los aspectos fundamentales de las redes de NGN.
3. Presentación de los elementos que formarán parte de la propuesta.
4. Elección del software de simulación.
5. Implementación de la propuesta en el software de simulación.
6. Análisis y discusión de los resultados obtenidos en la simulación.

Firma del Autor

Firma del Tutor

RESUMEN

La tendencia en el mundo actual de las Telecomunicaciones e Informática **permitirá** la convergencia de voz, video y datos en una sola red conmutada en paquetes mediante el predominio del protocolo IP y las redes NGN, las cuales facilitan el aumento de la velocidad en el acceso a los servicios de banda ancha, las Redes de Acceso Móviles de voz y datos. Para garantizar todas estas posibilidades es necesario tener un núcleo de red que sirva como soporte y que cumpla con las exigencias de las redes actuales. En este trabajo **realiza** una propuesta de núcleo MPLS que cumple con estos requisitos, **haciendo** marcado énfasis en cumplir con los requerimientos de calidad de servicio y garantizando la escalabilidad y estabilidad de la red. Para ello se realiza una caracterización de las principales tecnologías asociadas a MPLS, se describen los elementos que interactúan en una red MPLS y a partir del software OPNET Modeler se hace una **evaluación** del funcionamiento de la red.

INDÍCE

PENSAMIENTO	iv
DEDICATORIA	v
AGRADECIMIENTOS	vi
TAREA TÉCNICA.....	vii
RESUMEN	viii
INDÍCE	ix
INTRODUCCIÓN	13
CAPÍTULO I: FUNDAMENTOS DE LAS TECNOLOGÍAS DE TRANSPORTE EN REDES DE PRÓXIMA GENERACIÓN	15
1.1 Redes de Nueva Generación.	15
1.1.1 Capa de Servicios/Aplicación.	15
1.1.2 Capa de Transporte.	16
1.1.3 Capa de Control.	16
1.1.4 Capa de Acceso.....	17
1.2 MPLS.	17
1.2.1 Definición.	18
1.2.2 Ideas preconcebidas sobre MPLS.	19
1.2.3 Integración y migración hacia MPLS.	21
1.2.4 Antecedentes de ATM.	21
1.2.5 Diferencia significativa entre MPLS y las tecnologías WAN.	21
1.3 Control de la información en MPLS.	23
1.3.1 Funcionamiento global MPLS.	23
1.4 Principales Aplicaciones de MPLS.....	24

1.4.1	Ingeniería de tráfico.	24
1.4.2	Clases de servicio.....	26
1.4.3	Redes Privadas Virtuales	26
1.4.3.1	Ventajas de VPNs con MPLS.	28
1.4.3.2	VPN en MPLS.....	28
1.5	GMPLS.	29
1.5.1	Funciones principales de GMPLS.	31
1.5.2	Tipos de conmutación y jerarquía de enrutamiento.....	32
1.5.3	Modelo de enrutado y direccionamiento.	33
1.5.4	Extensiones del plano de control de MPLS.	33
1.5.5	Mejoras de escalabilidad en GMPLS.....	35
1.6	Conclusiones parciales del capítulo.	35
CAPÍTULO II: ELEMENTOS DEL NÚCLEO DE LA RED		36
2.1	Componentes de MPLS.....	36
2.1.1	LER.....	36
2.1.2	LSR.	37
2.1.3	LSP.....	37
2.1.4	FEC.	38
2.1.4.1	Agregación.	38
2.1.5	LDP.....	39
2.1.6	LIB.....	39
2.2	Funcionamiento del envío de paquetes en MPLS.	40
2.3	Mecanismos de encolamiento.	43
2.3.1	FIFO.....	43

2.3.2	WFQ.....	44
2.3.3	Class-Based WFQ.....	44
2.3.4	PQ.	45
2.3.5	CQ.....	46
2.4	Modos de simulación en OPNET.....	47
CAPÍTULO III: PROPUESTA DEL NÚCLEO NGN		49
3.1	Planteamiento del diseño.....	49
3.2	Diagrama de Topología.....	49
3.3	Aplicaciones y Perfiles.....	50
3.4	Configuración de las VPN en la red IP/MPLS.....	50
3.4.1	Creación de las interfaces de lazo cerrado.....	50
3.4.2	Creación de las VRF.....	51
3.4.3	Asignación de las VRF.....	51
3.4.4	Protocolo de enrutamiento OSPF.....	51
3.4.5	Border Gateway Protocol.....	52
3.4.6	Familia de Direcciones para IPv4 VRF.....	52
3.4.7	Calidad de Servicio.....	53
3.5	Evaluación de los resultados de la propuesta.....	53
CONCLUSIONES Y RECOMENDACIONES.....		63
REFERENCIAS BIBLIOGRÁFICAS.....		64
GLOSARIO		67
ANEXOS		70
Anexo 1:	Tráfico recibido para BGP.....	70
Anexo 2:	Tiempo de Respuesta para descarga y subida de archivos Email.....	70

Anexo 3:	Tráfico enviado y recibido para aplicaciones Email.....	71
Anexo 4:	Tiempo de Respuesta para descarga y subida de archivos FTP.	71
Anexo 5:	Tráfico recibido y enviado para la aplicación FTP.....	72
Anexo 6:	Tiempo de respuesta para un objeto y una página utilizando el protocolo HTTP.	72
Anexo 7:	Tráfico recibido y enviado para la aplicación HTTP.....	73
Anexo 8:	Tráfico recibido y enviado para aplicaciones de Voz.....	73
Anexo 9:	Razón de Transferencia para las demandas simulando el tráfico de videoconferencia.....	74
Anexo 10:	Razón de Transferencia punto a punto para tráfico de voz existente entre los softswitch y la red NGN Villa Clara.....	74
Anexo 11:	Razón de Transferencia punto a punto para tráfico de voz existente entre los softswitch Las Tunas y Vedado.	75
Anexo 12:	Perfiles creados en la Red.....	75
Anexo 13:	Tabla VRF de Audio para el ISAM de Santa Clara.....	76
Anexo 14:	Red simulada y modelada para Villa Clara.	76

INTRODUCCIÓN

Las redes de telecomunicaciones, como soporte para el transporte de información, han alcanzado un avance espectacular en las últimas décadas y han tenido un papel protagonista en la evolución de las tecnologías de la información y las comunicaciones.

La tendencia en el mundo actual de las Telecomunicaciones e Informática permitirá la convergencia de voz, video y datos en una sola red conmutada en paquetes mediante el predominio del protocolo IP y con él surgen las (NGN), las cuales facilitan el aumento de la velocidad en el acceso a los servicios de banda ancha, las Redes de Acceso Móviles de voz y datos.

La evolución del sector hacia las redes convergentes o (NGN) está ligada a la evolución del estado hacia la Sociedad de la Información, en la medida en que estas redes constituyen la principal infraestructura para el transporte de la información y para la conectividad de las personas.

Cuba por diversas razones no puede marchar a la par del mundo desarrollado en algunos aspectos de la tecnología, dígame costos de las mismas. En la medida que las posibilidades lo han permitido la tecnología con que se cuenta es cada vez más cercana a la de los países desarrollados. Por tanto un cambio de concepción en cuanto a arquitectura y estructura de redes se impone, en aras de lograr una convergencia e integración de todos los servicios de telecomunicaciones que se brindan en la actualidad.

Por lo que surge como necesidad la implementación de un *Backbone* en la provincia de Villa Clara para ofrecer a los usuarios seguridad en cuanto a temas como QoS y TE, por tan solo mencionar algunos.

Las redes de telecomunicaciones en sus inicios no fueron diseñadas para soportar las aplicaciones, de gran ancho de banda, existentes hoy en día, por lo que surgen las redes NGN, que no es más que una red basada en paquetes que permite prestar servicios de telecomunicación y en la que se pueden utilizar múltiples tecnologías de transporte de banda ancha propiciadas por la QoS, y en la que las funciones relacionadas con los servicios son independientes de las tecnologías subyacentes relacionadas con el transporte. Todo esto conlleva a la siguiente **situación problemática**: ¿cómo garantizar el correcto

funcionamiento del *Backbone* de la red NGN en Villa Clara a partir de una solución IP/MPLS?

Para ello se plantea el siguiente **objetivo general**: realizar una propuesta basada en soluciones IP/MPLS para el núcleo de la red NGN de Villa Clara que garantice QoS.

Los objetivos específicos son:

1. Describir los fundamentos de las tecnologías asociadas al núcleo de una red NGN.
2. Describir los elementos que forman parte de la propuesta.
3. Evaluar la propuesta mediante Simulación.

Con este proyecto se pretende contribuir al perfeccionamiento del *Backbone* de la red NGN de Villa Clara. Además esta investigación puede servir como material bibliográfico para futuros proyectos investigativos relacionados con el tema.

El informe de la investigación se estructurará en introducción, tres capítulos, conclusiones, recomendaciones, referencias bibliográficas, glosario y anexos. En la introducción se dejará definida la importancia, actualidad y necesidad del tema que se aborda y se dejarán explícitos los elementos del diseño teórico. En el capítulo 1 se describirán las principales características de las redes NGN haciendo énfasis en lo relacionado con las diferentes tecnologías de transporte. En el capítulo 2 se definirán los elementos **tener** en cuenta, las características de los mismos, y el software de simulación. En el capítulo 3 se realizará la propuesta y se expondrán los resultados obtenidos en las simulaciones para de esta forma evaluar el desempeño de una red NGN para servicios de tiempo real. A continuación las conclusiones estarán en función del cumplimiento de los objetivos y las recomendaciones en correspondencia a futuros trabajos. Las referencias bibliográficas revelan la novedad y actualidad del tema. A continuación se muestra el glosario y los anexos.

CAPÍTULO I: FUNDAMENTOS DE LAS TECNOLOGÍAS DE TRANSPORTE EN REDES DE PRÓXIMA GENERACIÓN

1.1 Redes de Nueva Generación.

El término (NGN), pese a las discrepancias de su precisión, o incluso, su corrección, se había arraigado a comienzos del año 2000 para designar las infraestructuras avanzadas del escenario convergente que se atisbaba, según la (ITU). Hoy día ya existen definiciones que describen qué pretenden las NGN y cuáles serán sus implementaciones. Una de las definiciones más acertadas la brinda la propia ITU.

Una Red de Próxima Generación es una red basada en la transmisión de paquetes capaz de proveer servicios integrados, incluyendo los tradicionales servicios telefónicos, capaz de explotar al máximo el ancho de banda del canal haciendo uso de las Tecnologías de (QoS) de modo que el transporte sea totalmente independiente de la infraestructura de red utilizada [1].

Esta tipología de red provee servicios de telecomunicaciones haciendo uso de múltiples tecnologías de banda ancha con transporte de QoS y con funciones relativas independientes. Pretende integrar las diferentes tecnologías presentes en los mercados actuales y las que se puedan desarrollar en el futuro, además de satisfacer las necesidades de información de los usuarios de forma transparente, de tal manera que los clientes no necesiten conocer la tecnología con que se atiende su demanda. La separación entre la porción de red de transporte (conectividad) y los servicios que corren por encima de esa red, permite que un proveedor telefónico que desee habilitar un nuevo servicio, pueda hacerlo fácilmente definiéndolo directamente desde la capa de servicio sin tener en cuenta la capa de transporte[1], [2].

1.1.1 Capa de Servicios/Aplicación.

En esta capa se encuentran algunos de los servicios y aplicaciones disponibles en una red NGN, los cuales son ofrecidos a toda la red sin importar la ubicación del usuario y son independientes de la tecnología de acceso utilizada.

Algunos de los servicios que se brindan a los usuarios son:

- Redes Inteligentes.
- Video en demanda.
- Correo electrónico.
- Correo de voz.

En esta capa se realizan las funciones relacionadas con la operación y administración de la red y sus servicios. Dentro de las tareas que se desarrollan en ella se encuentran la gestión de fallas, configuración de red y elementos, medición del desempeño, **tasación**, seguridad, gestión de tráfico y QoS.

Además esta capa la componen los servidores de aplicaciones y de medios, los cuales son los encargados de proveer las funciones y características de la red. Dentro de estas características está el establecimiento de las conexiones, el encaminamiento, la facturación y los servicios avanzados que son posibles de implementar por medio de la señalización y la información que se deduce de esta capa [2], [3].

1.1.2 Capa de Transporte.

Esta capa **no es más** que el *backbone* (red dorsal) de alta velocidad, de transmisión óptica, el cual soportará el tráfico de paquetes para todos los servicios, es decir, voz, datos, video y otros, es responsable de la QoS de extremo a extremo. Mantendrá conectividad entre todos los componentes y la separación física entre las funciones dentro de NGN. El equipamiento que lo compone son enrutadores y conmutadores que permitirán la conmutación de las señales por la red asegurando alta capacidad y confiabilidad. Este nivel adopta tecnología de conmutación de paquetes IP, o ATM, pero el IP se reconoce como la tecnología de transporte más prometedora para NGN [4].

1.1.3 Capa de Control.

La capa de control es la más importante dentro de la arquitectura de la NGN. En esta capa se encuentra los dispositivos que controlan todo el transporte de los datos en la red así como el acceso a la misma. Esta capa proporciona las funciones de control de los servicios y de la red. Esta se encarga de asegurar el interfuncionamiento de la red de transporte con

los servicios y aplicaciones, mediante la interpretación, generación, distribución, y traducción de la señalización correspondiente, mediante el uso de diferentes protocolos. La separación del control y la inteligencia de la red de las funciones de transporte es una característica intrínseca del diseño de la NGN. Comprende las funciones de señalización y ruteo de las llamadas y conexiones en toda la red. Los dispositivos y funciones en esta capa llevan a cabo el control basado en mensajes de señalización recibidos de la capa de transporte y el establecimiento y terminación de conexiones multimedia a través de la red, controlando también componentes en la capa de transporte.

Una de las arquitecturas más comunes en NGN es el *softswitch*. En esta capa es donde se encuentra el *softswitch* que define dicha arquitectura, es el dispositivo de control principal de la red, mediante la combinación de *hardware* y *software* [5].

1.1.4 Capa de Acceso.

Esta capa incluye una diversidad de tecnologías usadas para llegar al cliente. Está compuesta por una variedad de dispositivos que permiten a los usuarios finales tener conectividad con la NGN, estos pueden ser MGs, AMGs TMGs o SMGs IADs y puntos de acceso inalámbrico. Además, existen los dispositivos que realizan las tres funciones de los tres primeros mencionados, estos son conocidos como UMG [2], [6].

1.2 MPLS.

En la década de los 90 varias empresas desarrollaron sus propias tecnologías para mezclar la alta velocidad de operación de ATM basada en conmutación, con el proceso de enrutamiento IP de Internet de la capa de red .

(CRS): Fue desarrollado por Toshiba y presentado al (IETF) en 1994. Fue una de las primeras propuestas para utilizar protocolos IP para controlar infraestructura ATM. CRS se ha desarrollado en redes comerciales y académicas en Japón.

IP Switching: Desarrollado por *Ipsilon Network* (ahora parte de Nokia), se anunció en 1996. El objetivo básico de *IP switching* fue el de integrar conmutadores ATM de una manera eficiente (eliminando el plano de control ATM). *Ipsilon* utilizó la presencia de tráfico para controlar el establecimiento de una etiqueta.

Tag Switching: Es la tecnología de conmutación de etiquetas desarrollada por *Cisco Systems*. A diferencia de las dos soluciones anteriores, *tag switching* es una técnica la cual no requiere de flujo de tráfico para la creación de tablas de etiqueta en un enrutador. En lugar de esto utilizaba protocolos de enrutamiento IP para determinar el siguiente salto.

(ARIS): Es desarrollado por IBM y es muy similar a *Tag switching* de Cisco. En ARIS la distribución de etiquetas comienza en el enrutador de salida y se propaga de forma ordenada hacia el enrutador de entrada.

La existencia de múltiples soluciones independientes para el desarrollo de tecnologías basadas en conmutación de etiquetas dio lugar a la necesidad de desarrollar estándares, fue en marzo de 1997 que IETF aprobó el grupo de trabajo MPLS, con el propósito de definir una propuesta basada en las normas para la conmutación de etiquetas [7], [8]. La evolución hasta llegar al estándar MPLS del IETF se ilustra en la siguiente:

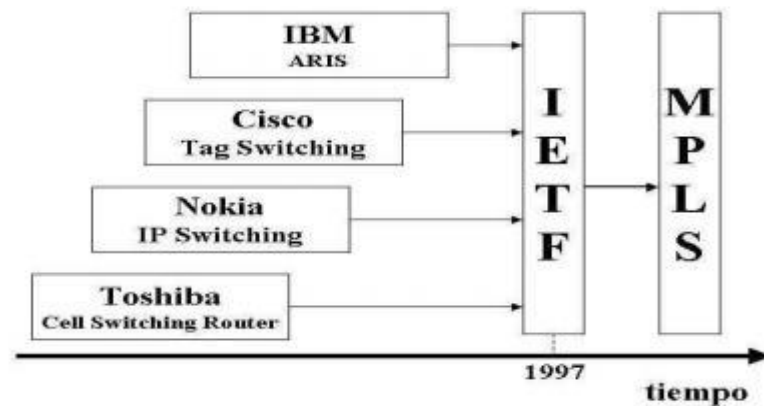


Fig.1.1: Evolución hacia MPLS [7].

1.2.1 Definición.

MPLS es un método para transmitir paquetes a través de una red usando información contenida en etiquetas añadidas a los paquetes de IP. MPLS se basa en el etiquetado de los paquetes en base a criterios de prioridad y/o QoS. La idea de MPLS es realizar la conmutación de los paquetes o datagramas en función de las etiquetas añadidas en la capa 2 y etiquetar dichos paquetes según la clasificación establecida por la QoS. Por tanto MPLS es una tecnología que permite ofrecer QoS, independientemente de la red sobre la que se

implemente. El etiquetado en la capa 2 permite ofrecer servicio multiprotocolo y ser portable sobre multitud de tecnologías de capa de enlace: ATM, Frame Relay, Ethernet, líneas dedicadas, (LANs, *Local Área Network*), etc. Debido a esto MPLS ha sido sin duda una de las tecnologías de más éxito de la última década [4], [7], [9].

1.2.2 Ideas preconcebidas sobre MPLS.

Durante el tiempo en que se ha desarrollado el estándar, se han extendido algunas ideas falsas o inexactas sobre el alcance y objetivos de MPLS. Hay quien piensa que MPLS se ha desarrollado para ofrecer un estándar a los vendedores que les permitiese evolucionar los conmutadores ATM a *routers* de *backbone* de altas prestaciones. Aunque esta puede haber sido la finalidad original de los desarrollos de conmutación multinivel, los recientes avances en tecnologías de silicio ASIC permiten a los routers funcionar con una rapidez similar para la consulta de tablas a las de los conmutadores ATM [10]. Si bien es cierto que MPLS mejora notablemente el rendimiento del mecanismo de envío de paquetes, éste no era el principal objetivo del grupo del IETF. Los objetivos establecidos por ese grupo en la elaboración del estándar eran:

- MPLS debía funcionar sobre cualquier tecnología de transporte, no sólo ATM.
- MPLS debía soportar el envío de paquetes tanto unicast como multicast.
- MPLS debía ser compatible con el Modelo de Servicios Integrados del IETF, incluyendo el (RSVP).
- MPLS debía permitir el crecimiento constante de la Internet.
- MPLS debía ser compatible con los procedimientos de operación, administración y mantenimiento de las actuales redes IP.

También ha habido quien pensó que el MPLS perseguía eliminar totalmente el encaminamiento convencional por prefijos de red. Esta es otra idea falsa y nunca se planteó como objetivo del grupo, ya que el encaminamiento tradicional de nivel 3 siempre sería un requisito en la Internet por los siguientes motivos:

- El filtrado de paquetes en los cortafuegos de acceso a las LAN corporativas y en los límites de las redes de los NSP es un requisito fundamental para poder gestionar la red y los servicios con las necesarias garantías de seguridad. Para ello se requiere examinar

la información de la cabecera de los paquetes, lo que impide prescindir del uso del nivel 3 en ese tipo de aplicaciones.

- No es probable que los sistemas finales implementen MPLS. Necesitan enviar los paquetes a un primer dispositivo de red (nivel 3) que pueda examinar la cabecera del paquete para tomar luego las correspondientes decisiones sobre su envío hasta su destino final. En este primer salto se puede decidir enviarlo por ruteo convencional o asignar una etiqueta y enviarlo por un *LSP*
- Las etiquetas MPLS tienen solamente significado local (es imposible mantener vínculos globales entre etiquetas y *hosts* en toda la Internet). Esto implica que en algún punto del camino algún dispositivo de nivel 3 debe examinar la cabecera del paquete para determinar con exactitud por dónde lo envía: por *routing* convencional o entregándolo a un *LSR* que lo expedirá por un nuevo *LSP*.
- Del mismo modo, el último *LSR* de un *LSP* debe usar encaminamiento de nivel 3 para entregar el paquete al destino, una vez suprimida la etiqueta, como se verá seguidamente al describir la funcionalidad MPLS [4], [11], [12].

El aprovechamiento de las funciones de conmutación del nivel 2 con el enrutamiento del nivel 3 ha provocado que MPLS se convierta en una arquitectura con gran porvenir. MPLS es una tecnología sencilla y barata, que puede acelerar el encaminamiento y la conmutación en los dorsales de las redes de datos con protocolos IP. MPLS proporciona un grupo de beneficios adicionales, por ejemplo:

- Es una solución que permite una gran interoperabilidad, no hay necesidad del modelo de superposición de IP sobre ATM y por tanto se eliminan las dificultades asociadas con su administración y la gestión.
- Proporciona los medios para el mapeo de direcciones IP a simples etiquetas de longitud fija utilizadas por diferentes tecnologías de envío y conmutación de paquetes.
- Permite el uso de diferentes protocolos existentes como son el RSVP, y OSPF.
- Permite proporcionar Diff-Serv, así como QoS e Ingeniería de tráfico y crear VPN, [9], [13].

1.2.3 Integración y migración hacia MPLS.

El estándar MPLS de la IETF pretende integrarse operando tanto con las tecnologías de transmisión de los niveles inferiores existentes, dígase, ATM, *Frame Relay*, PPP SDH, y LANs, como con los diferentes protocolos de nivel de red, dígase, IP e IPX, Las redes MPLS se han ido integrando y formando las redes actuales, a partir de dos infraestructuras fundamentales, ATM y redes basadas en enrutadores IP puros [7].

1.2.4 Antecedentes de ATM.

El Nodo ATM es una tecnología que ofrece un servicio orientado a la conexión y trabaja en el nivel 2 del modelo OSI, esta tecnología se basa en estándares y es utilizada para proporcionar servicios multimedia, utilizando conmutación de celdas de longitud fija y circuitos virtuales para el transporte de datos, voz y video de manera rápida. ATM combina los beneficios de la conmutación de circuitos (retardo constante de transmisión y capacidad garantizada) con los beneficios de la conmutación de paquetes (flexibilidad y eficiencia en el uso del ancho de banda) [10].

Con MPLS los enlaces ATM son tratados como enlaces IP y cada conmutador se convierte además en un enrutador IP modelo integrado. Implementando la inteligencia IP en los conmutadores ATM, se pueden diseñar redes en la que se eliminen las superposiciones de enlaces IP sobre ATM, logrando una correspondencia uno a uno entre ellos, lo cual resuelve la mayoría de los problemas de escalabilidad IP. Por otra parte con la integración de los Niveles 2 y 3 se alcanza un modelo que toma las mejores características de cada nivel.

Un concepto básico de MPLS es convertir el equipo de Nivel 2 en la red de conmutadores ATM en ATM-LSRs, el cual puede ser visto como una combinación de un sistema de conmutación ATM y un enrutador tradicional, compuesto por el plano de control y el plano de envío [14].

1.2.5 Diferencia significativa entre MPLS y las tecnologías WAN.

La diferencia significativa entre MPLS y las tecnologías WAN tradicionales es la forma en que son asignadas las etiquetas y la capacidad de portar una pila de etiquetas adjuntadas al

paquete. El concepto de pila de etiquetas posibilita nuevas aplicaciones como (TE), (VPN), **re enrutamiento** rápido ante fallos ante enlaces o nodos y otras.

El uso de MPLS, tecnología orientada a conexión a través de establecimientos de trayectorias en el transporte de paquetes, difiere de las redes no orientadas a conexión de hoy, donde cada paquete es analizado básicamente salto a salto, es chequeada su cabecera de Nivel 3, y de una forma independiente se van tomando decisiones de envío basadas en información extraída de los algoritmos de enrutamiento del nivel de red.

La arquitectura MPLS es dividida en dos componentes separados: la componente de envío, también llamada plano de datos y la componente de control, conocida como plano de control. La componente de envío usa una base de datos de las etiquetas enviadas, que es mantenida por el conmutador de etiquetas para ejecutar el envío de paquetes portando etiquetas. El componente de control esta responsabilizado con la creación y mantenimiento de la información de las etiquetas enviadas en medio de un grupo de conmutadores de etiquetas interconectados [4], [10], [15].

La Figura 1.2 muestra la arquitectura básica de un nodo MPLS ejecutando un enrutamiento IP.

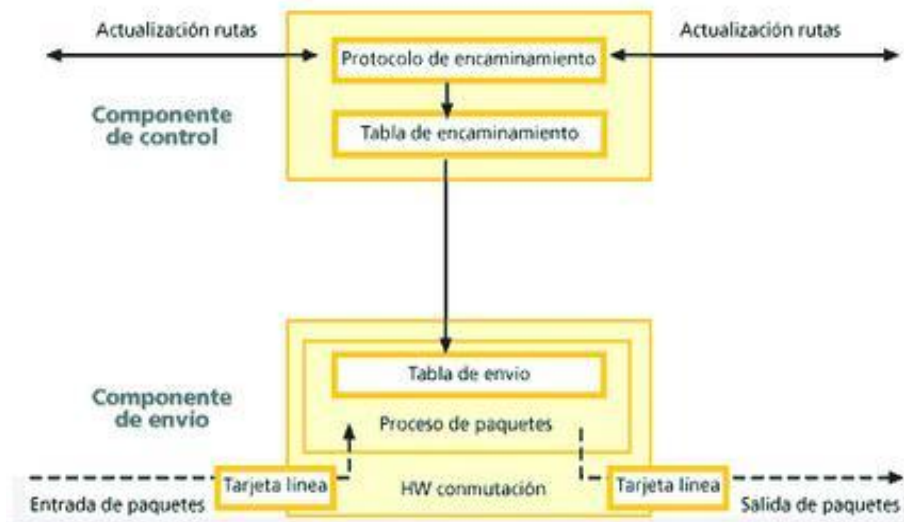


Figura.1.2: Arquitectura Básica de un Nodo MPLS [4].

1.3 Control de la información en MPLS.

El primero de ellos está relacionado con la información que se tiene sobre la red: topología, patrón de tráfico, características de los enlaces, etc. Es la información de control típica de los algoritmos de encaminamiento. MPLS necesita esta información de *routing* para establecer los LSPs. Lo más lógico es utilizar la propia información de encaminamiento que manejan los protocolos internos IGP, OSPF, IS-IS, RIP para construir las tablas de encaminamiento. Esto es lo que hace MPLS precisamente para cada "ruta IP", en la red se crea un "camino de etiquetas" a base de concatenar las de entrada/salida en cada tabla de los LSRs; el protocolo interno correspondiente se encarga de pasar la información necesaria.

Siempre que se quiera establecer un circuito virtual se necesita algún tipo de señalización para marcar el camino, es decir, para la distribución de etiquetas entre los nodos. Sin embargo, la arquitectura MPLS no asume un único protocolo de distribución de etiquetas; de hecho se están estandarizando algunos existentes con las correspondientes extensiones; uno de ellos es el protocolo RSVP del Modelo de Servicios Integrados del IETF. Pero, además, en el IETF se están definiendo otros nuevos, específicos para la distribución de etiquetas, cual es el caso del LDP,[15], [16].

1.3.1 Funcionamiento global MPLS.

Una vez vistos todos los componentes funcionales, el esquema global de funcionamiento es el que se muestra en la figura 1.4, donde quedan reflejadas las diversas funciones en cada uno de los elementos que integran la red MPLS. Es importante destacar que en el borde de la nube MPLS existe una red convencional de routers IP. El núcleo MPLS proporciona una arquitectura de transporte que hace aparecer a cada par de routers a una distancia de un sólo salto. Funcionalmente es como si estuvieran unidos todos en una topología mallada. Ahora, esa unión a un solo salto se realiza por MPLS mediante los correspondientes LSPs (puede haber más de uno para cada par de routers). La diferencia con topologías conectivas reales es que en MPLS la construcción de caminos virtuales es mucho más flexible y que no se pierde la visibilidad sobre los paquetes IP. Todo ello abre enormes posibilidades a la hora de mejorar el rendimiento de las redes y de soportar nuevas aplicaciones de usuario [4].

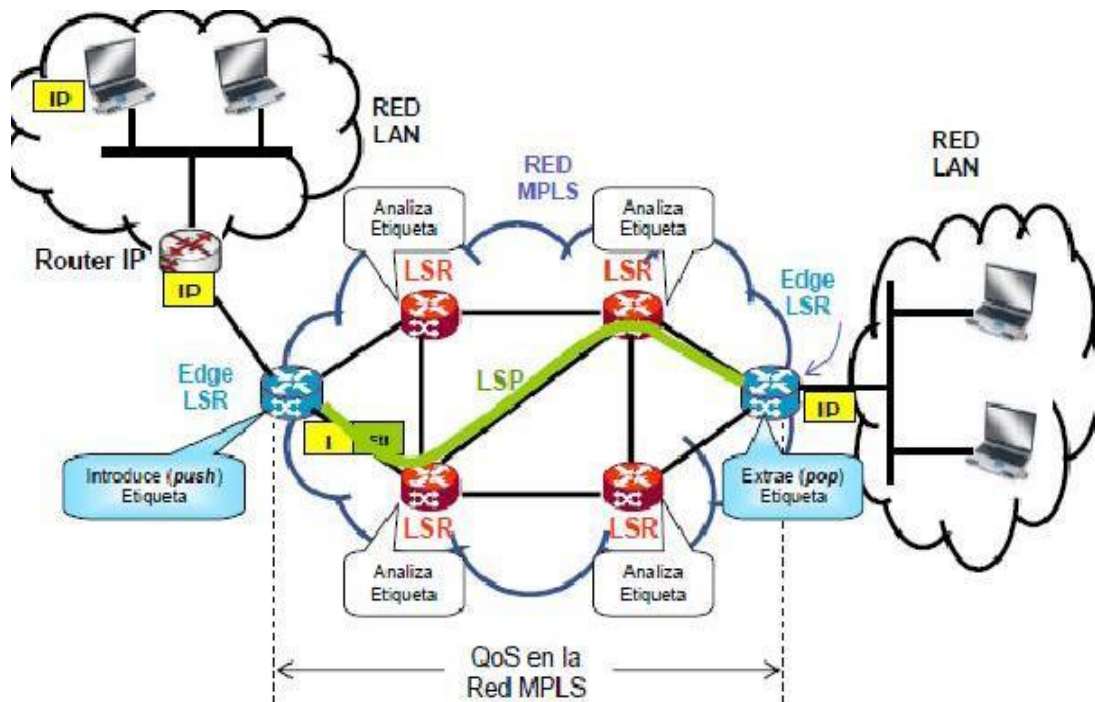


Fig: 1.4 Esquema global de funcionamiento que integran la red MPLS [4].

1.4 Principales Aplicaciones de MPLS.

1.4.1 Ingeniería de tráfico.

El objetivo básico de la ingeniería de tráfico es adaptar los flujos de tráfico a los recursos físicos de la red. La idea es equilibrar de forma óptima la utilización de esos recursos, de manera que no haya algunos que estén supra utilizados, con posibles puntos calientes y cuellos de botella, mientras otros puedan estar infrautilizados. A comienzos de los 90 los esquemas para adaptar de forma efectiva los flujos de tráfico a la topología física de las redes IP eran bastante rudimentarios. Los flujos de tráfico siguen el camino más corto calculado por el algoritmo IGP correspondiente. En casos de congestión de algunos enlaces, el problema se resolvía a base de añadir más capacidad a los enlaces. La ingeniería de tráfico consiste en trasladar determinados flujos seleccionados por el algoritmo IGP sobre enlaces más congestionados, a otros enlaces más descargados, aunque estén fuera de la ruta más corta (con menos saltos). En la figura 1.5 se comparan estos dos tipos de rutas para el mismo par de nodos origen-destino [17]–[19].

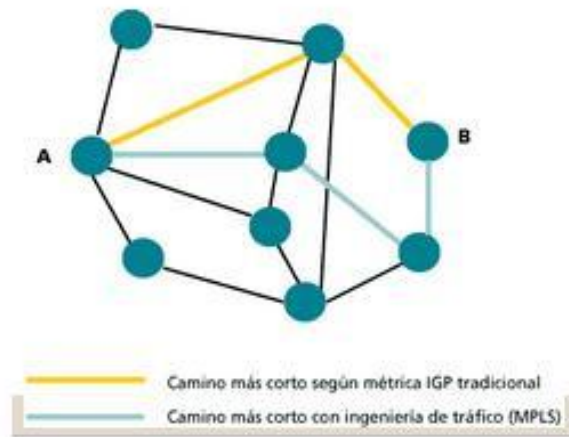


Fig. 1.9: Comparación de dos tipos de rutas [19].

El camino más corto entre A y B según la métrica normal IGP es el que tiene **sólo** dos saltos, pero puede que el exceso de tráfico sobre esos enlaces o el esfuerzo de los routers correspondientes, hagan aconsejable la utilización del camino alternativo indicado con un salto más. MPLS es una herramienta efectiva para esta aplicación en grandes **backbones**, ya que:

- Permite al administrador de la red el establecimiento de rutas explícitas, especificando el camino físico exacto de un LSP.
- Permite obtener estadísticas de uso LSP, que se pueden utilizar en la planificación de la red y como herramientas de análisis de cuellos de botella y carga de los enlaces, lo que resulta bastante útil para planes de expansión futura.
- Permite hacer encaminamiento restringido CBR de modo que el administrador de la red pueda seleccionar determinadas rutas para servicios especiales (distintos niveles de calidad). Por ejemplo, con garantías explícitas de retardo, ancho de banda, fluctuación, pérdida de paquetes, etc.

La ventaja de la ingeniería de tráfico MPLS es que se puede hacer directamente sobre una red IP, al margen de que haya o no una infraestructura ATM por debajo, todo ello de manera más flexible y con menores costes de planificación y gestión para el administrador, y con mayor QoS para los clientes [8], [20].

1.4.2 Clases de servicio.

MPLS está diseñado para poder cursar servicios diferenciados, según el Modelo *DiffServ* del IETF. Este modelo define una variedad de mecanismos para poder clasificar el tráfico en un reducido número de CoS, con diferentes prioridades. Según los requisitos de los usuarios, *DiffServ* permite diferenciar servicios tradicionales tales como el **www**, el correo electrónico o la transferencia de ficheros (para los que el retardo no es crítico), de otras aplicaciones mucho más dependientes del retardo y de la variación del mismo, como son las de vídeo y voz interactiva. Para ello se emplea el campo ToS en el modelo *DiffServ* se le conoce como el octeto DS. Esta es la técnica QoS de marcar los paquetes que se envían a la red [13], [21], [22].

MPLS se adapta perfectamente a ese modelo, ya que las etiquetas MPLS tienen el campo EXP, para poder propagar la CoS en el correspondiente LSP. De este modo, una red MPLS puede transportar distintas ToS, ya que:

- El tráfico que fluye a través de un determinado LSP se puede asignar a diferentes colas de salida en los diferentes saltos LSR, de acuerdo con la información contenida en los bits del campo EXP.
- Entre cada par de LSR exteriores se pueden provisionar múltiples LSPs, cada uno de ellos con distintas prestaciones y con diferentes garantías de ancho de banda. Los LSP brindan, tres niveles de servicio, primero, preferente y turista, que, lógicamente, tendrán distintos precios.

1.4.3 Redes Privadas Virtuales

Una VPN se construye a base de conexiones realizadas sobre una infraestructura compartida, con funcionalidades de red y de seguridad equivalentes a las que se obtienen con una red privada. El objetivo de las VPNs es el soporte de aplicaciones intra/extranet, integrando aplicaciones multimedia de voz, datos y vídeo sobre infraestructuras de comunicaciones eficaces y rentables. La seguridad supone aislamiento, y "privada" indica que el usuario "cree" que posee los enlaces. Las IP VPNs son soluciones de comunicación VPN basada en el protocolo de red IP de la Internet.

Las VPNs tradicionales ya sean basadas en PVC o túneles IP han sido de gran beneficio pero tienen ciertos inconvenientes que pueden ser resueltos con la utilización de MPLS. Las VPNs basadas en PVC utilizan la infraestructura de las redes ATM o *Frame Relay* y los PVCs se establecen entre los nodos de extremo a extremo con la configuración manual de cada uno, lo que implica complejidad en la gestión de la red del proveedor ya que se trata de una topología lógica mallada sobrepuesta a la red física y al agregar un nuevo miembro a la VPN es necesario restablecer todos los PVCs, la seguridad que ofrecen se basa en la separación de tráfico por PVC [23].

Las IP VPN están basadas en Protocolos de Túnel como por ejemplo *IPSec*, la información se cifra y se encapsula en una nueva cabecera IP. La desventaja en este tipo de implementaciones se da porque se ocultan las cabeceras de los paquetes originales y las opciones de QoS son bastante limitadas ya que no se puede distinguir los flujos por aplicación dificultando la asignación de los diferentes niveles de servicio [24].

En general los inconvenientes más comunes que tienen las VPN tradicionales son las siguientes:

- Se basan en conexiones punto a punto (PVC o túneles).
- La configuración de cada nodo de la VPN es manual y cada vez que se integra uno supone la reconfiguración de todos los anteriores.
- La QoS se ofrece hasta cierta parte, más no durante el transporte.
- El modelo topológico sobrepuesto a la red existente implica poca flexibilidad en la provisión y gestión del servicio.

Utilizando MPLS para implementar VPNs se eliminan los inconvenientes de las tecnologías anteriores. En primera instancia el modelo topológico que se crea no se sobrepone sino se acopla a la red del proveedor, esto elimina las conexiones extremo a extremo (túneles IP convencionales o circuitos virtuales) y los túneles se van creando con el intercambio de las etiquetas formándose así los LSP que vendrían a ser los “túneles MPLS”. Dentro de la red del proveedor las VPNs se forman mediante las rutas virtuales LSPs, similares a los túneles de las VPNs tradicionales pero con la diferencia de que la información se transporta por el mecanismo de intercambio de etiquetas obviando la

información de enrutamiento lo que facilita aplicar técnicas de QoS que son propagadas hasta el destino, reservando ancho de banda, estableciendo CoQ y aplicando TE de esta manera optimizando los recursos de la red y cumpliendo los máximos requerimientos de disponibilidad y seguridad [7], [25].

1.4.3.1 Ventajas de VPNs con MPLS.

- Se elimina la complejidad de los túneles y los PVCs.
- Para la implementación no es necesario realizar cambios en todos los puntos involucrados como ocurre con las VPNs tradicionales por lo contrario solo se configura a nivel del proveedor evitando tareas complejas y riesgosas.
- Las garantías de QoS se mantienen de extremo a extremo separando los flujos de tráfico por clases.
- Para aumentar la seguridad se pueden utilizar los protocolos de encriptación manejados también por las VPNs tradicionales como IPSec.
- Con la TE que ofrece MPLS se garantiza que en el servicio VPN no influyan parámetros que afecten la calidad de extremo a extremo [26], [27].

1.4.3.2 VPN en MPLS.

Una VPN/MPLS básica está formada de tres elementos físicos que son: *Provider*, PE o router frontera del proveedor y CE denominado así al router frontera del cliente; estos elementos constituyen la arquitectura externa de una VPN. Además existen dos aspectos internos de la arquitectura de las VPN soportadas en MPLS que es necesario mencionar para una mejor comprensión de su funcionamiento y son: el *Route Distinguisher* y el *Route Target*, mecanismos que permiten distinguir los requerimientos del cliente suscrito a una VPN [23], [27].

1. **Route Distinguisher:** Los routers PE se conectan a los routers CE y distribuyen la información que contienen sobre las VPNs a otros routers PE a través del protocolo MP-BGP o Multiprotocolo BGP, en este intercambio de información el router PE agrega como prefijo a la dirección IPv4 una cantidad de 64 bits conocidos como *Route Distinguisher* lo que permite a la dirección IPv4 hacerla globalmente única (ruta privada) y resultando finalmente una dirección de 96 bits denominada VPNv4.

2. **Route Target:** El *Route Target* o ruta objetivo es un atributo adicional colocado a las rutas VPNv4 vía BGP que permite identificar la membresía de un cliente a una VPN cuando algunos sitios de cliente participan en más de una VPN. El *Route Target* se introdujo en la arquitectura de las VPN/MPLS para soportar topologías más complejas.

Cada enrutador PE define un valor numérico llamado *Route Target* que pueden ser:

- **Export:** es un valor numérico que identifica una membresía de VPN y es adjuntada a la ruta del cliente cuando se convierte en una ruta VPNv4.
- **Import:** cuando el valor numérico de una VPNv4 de la red de origen (*Route Target Export*) coincide con los valores del *Route Target Import* del route PE de destino, esta ruta es insertada en la tabla de ruteo virtual del router PE correspondiente al sitio del cliente de destino.

Para que una nueva ruta sea aceptada el valor de su ruta objetivo de salida (exportación) debe de coincidir con el valor de entrada (importación) del dispositivo de entrada [4], [25].

1.5 GMPLS.

La Conmutación de GMPLS corresponde al siguiente desarrollo evolutivo desde (MPLS-TE) encaminado a proporcionar características de redes orientadas a entornos no orientadas a conexión. Debido a la necesidad de asignar el tráfico IP directamente sobre la capa óptica para reducir la complejidad de las conexiones y permitir la rápida asignación del ancho de banda y flexibilidad IP, el grupo de trabajo de MPLS continuó sus investigaciones y en octubre de 2004 publican la **RFC-3945** en la que se describe GMPLS extiende MPLS para abarcar división de tiempo, SONET SDH PDH longitud de onda y la conmutación en el espacio.

GMPLS abarca, además de los enrutadores IP y los conmutadores ATM, dispositivos como, conmutadores de longitudes de onda con conversión electroóptica, sistemas DWDM, interconexiones fotónicas, etc. Para ello, GMPLS extiende ciertas funciones con base del tradicional MPLS y, en algunos casos, añade nueva funcionalidad. Estas adaptaciones han supuesto la extensión de los mecanismos de etiqueta y de LSP para caminos de etiquetas generalizados (G-LSP) afectando también **los protocolos** de encaminamiento y

señalización para actividades tales como la distribución de etiquetas, la TE, la protección y restauración de enlaces [28].

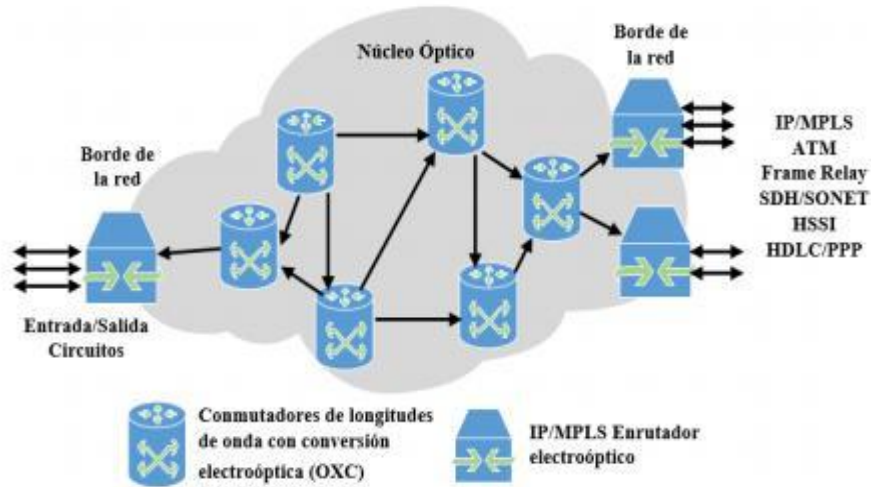


Fig.1.6: Arquitectura GMPLS [28].

El enfoque de la arquitectura GMPLS, mostrado en la figura 1.6, ejemplifica su primera diferencia en la manera de cubrir tanto la señalización como el enrutamiento desde un plano de control totalmente separado del plano de datos, tal como se muestra en la figura 1.7, a través del uso de unas reglas de gestión de enlace. Estos planos de control están compuestos de enrutadores MPLS-TP.

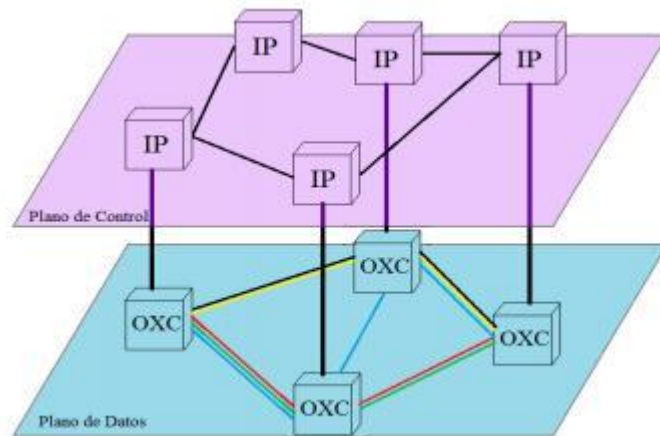


Fig.1.7: Plano de Control y Plano de Datos de la arquitectura GMPLS [28].

1.5.1 Funciones principales de GMPLS.

- Descubrimiento de vecinos: Con el fin de gestionar el núcleo de la red, cada componente conectado debe ser reconocido. Entre otros dispositivos se incluye conmutadores, enrutadores y multiplexores.
- Difusión de estado del enlace: En GMPLS, los protocolos de enrutamiento OSPF e IS - IS son modificados para tal fin, añadiéndose vínculos de tipo TE.
- Gestión de rutas: El protocolo CR-LDP y RSVP-TE se encargan del proceso de señalización.
- Gestión de enlaces: Necesaria para desarrollar y lanzar un total de canales ópticos con el fin de aumentar la escalabilidad.
- Protección y recuperación: Se puede crear una estructura malla que permita la aparición de un número de caminos diferentes.

GMPLS ofrece un panel de control único e integrado y extiende la disponibilidad de recursos y gestión del ancho de banda a lo largo de todas las capas de la red; e otras palabras, los equipos dejan de estar separados en diferentes capas y todos los elementos pueden tener información sobre el resto. Debido a que está diseñado para soportar diferentes tipos de tráfico, las redes presentan escalabilidad y simplicidad mediante el desarrollo de una nueva clase de elemento diseñado para manejar diferentes tipos de flujos simultáneamente [29], [30].

El principal beneficio que GMPLS ofrece actualmente a los operadores es una rápida provisión de servicios de cualquier tipo, con garantías de QoS, con cualquier grado de disponibilidad. Esta provisión tiene además un coste operativo muy bajo, por utilizar las ampliamente disponibles herramientas de gestión IP y apoyarse en un plano de control idéntico para gestionar la estructura óptica. Pero GMPLS permite también evolucionar gradualmente de una compleja red de transporte de datos de varias capas, a un modelo formado únicamente por dos. Esto es debido a que la funcionalidad proporcionada actualmente por las capas ATM y SONET/SDH, como la TE, la QoS, o las VPN serán progresivamente proporcionadas por las redes fotónicas DWDM mediante GMPLS [31], [32].

1.5.2 Tipos de conmutación y jerarquía de enrutamiento.

GMPLS soporta múltiples tipos de conmutación: conmutación TDM longitudes de onda y por fibra. Para incluir estos tipos adicionales de conmutación se han extendido ciertas funciones base de MPLS. Estos cambios y añadidos afectan a las propiedades básicas de los LSP, a cómo se solicitan y comunican las etiquetas, a la naturaleza unidireccional de los LSPs, a cómo se propagan los errores y a la información proporcionada para sincronizar los LSR de entrada y salida.

La arquitectura original de MPLS se ha extendido para incluir un nuevo conjunto de interfaces en los LSR. Estas interfaces se clasifican en:

- Interfaces PSC: Estas interfaces reconocen los límites de paquetes y pueden enviar datos basándose en el contenido de la cabecera de paquete.
- Interfaces L2SC: Estas interfaces reconocen los límites de tramas/celdas y pueden enviar datos basándose en el contenido de la cabecera de las tramas/celdas.
- Interfaces TDM: Estas interfaces enrutan los datos basándose en la ranura temporal de los datos dentro de un ciclo de repetición.
- Interfaces LSC: Estas interfaces enrutan los datos basándose en la longitud de onda sobre la que se reciben los datos.
- Interfaces FSC: Estas interfaces enrutan los datos basándose en la posición en que se reciben éstos en el espacio físico.

Un circuito sólo puede ser establecido entre interfaces del mismo tipo. Genéricamente todos los distintos tipos de circuitos que se pueden establecer entre dos interfaces del mismo tipo reciben el nombre de LSPs.

Un LSP puede anidarse dentro de otro creándose una jerarquía de LSPs. Para hacer esto se considera el LSP como un enlace en la base de datos de OSPF o IS-IS.

Los LSPs que entran y abandonan el dominio del transporte óptico en el mismo nodo pueden agregarse y ser encapsulados en un único LSP. Esta agregación permite conservar el número de longitudes de ondas que se utilizan en el dominio MPLS [32].

1.5.3 Modelo de enrutado y direccionamiento.

GMPLS se basa en los modelos de enrutado y direccionamiento IP. Esto implica que las direcciones IPv4 e IPv6 se utilizan para identificar las interfaces y que los protocolos de enrutamiento distribuidos tradicionales también se utilizan. El descubrimiento de la topología y el estado de los recursos de todos los enlaces en un dominio de enrutado se consigue a través de estos protocolos.

Dado que el plano de control y el de datos están desacoplados en GMPLS, un vecino del plano de control no tiene porqué serlo del plano de datos, por ello se requieren mecanismos como el LMP para asociar los enlaces TE con los nodos vecinos. Las direcciones IP no se utilizan únicamente para identificar interfaces en hosts IP y routers. De manera general, identifican cualquier interfaz PSC y no PSC. De igual manera, los protocolos de enrutamiento IP se utilizan para encontrar rutas para los datagramas IP con un algoritmo SPF. En un modelo superpuesto, cada capa no PSC en particular, se puede ver como un conjunto de AS interconectados de manera arbitraria. Análogamente al enrutado tradicional IP, cada AS es gestionado por una única autoridad administrativa. El intercambio de la información de enrutado entre entidades AS puede realizarse a través de un protocolo de enrutamiento entre dominios como BGP-4. Cada AS puede estar subdividido en distintos dominios de enrutamiento y cada uno puede ejecutar un protocolo de enrutamiento intra-dominio.

Cada dominio de enrutamiento puede estar dividido en áreas. Un RD se compone de nodos habilitados para GMPLS. Estos nodos pueden ser nodos de frontera (*Edge*) o LSRs internos. GMPLS define extensiones a los protocolos de enrutamiento intra-dominio OSPF e IS-IS. Estas extensiones son necesarias para diseminar características estáticas y dinámicas relacionadas con los nodos y los enlaces [31], [33].

1.5.4 Extensiones del plano de control de MPLS.

GMPLS extiende los planos de control originales de MPLS y/o MPLS-TE para soportar cada una de las cinco interfaces definidas anteriormente. El plano de control de GMPLS se compone de protocolos de señalización y enrutamiento conocidos que han sido modificados para soportar GMPLS. Estos protocolos utilizan direcciones IPv4 y/o IPv6. Sólo se necesita

un protocolo especializado para soportar las operaciones de GMPLS: un protocolo para la gestión de enlaces, LMP.

LMP proporciona mecanismos para mantener la conectividad del canal de control, verificar la conectividad física de los enlaces de datos, correlacionar la información de propiedad del enlace y gestionar los fallos en los enlaces.

LMP está definido en el contexto de GMPLS, pero está especificado independientemente de la señalización GMPLS, por lo que puede utilizarse en otros contextos y con otros protocolos de señalización.

Los protocolos de señalización y enrutamiento de MPLS requieren al menos un canal de control bidireccional para comunicar incluso si dos nodos adyacentes están conectados mediante enlaces unidireccionales. LMP puede establecer, mantener y gestionar estos canales de control. Los canales de control pueden transportarse «en banda» o «fuera de banda».

La mayoría de tecnologías que se pueden utilizar por debajo del nivel PSC requieren cierto nivel de TE. El establecimiento de LSPs a estos niveles, se deben tener en cuenta ciertas restricciones que no permiten utilizar el algoritmo SPF. En su lugar se utiliza enrutamiento SPF basado en restricciones. Los nodos que establecen los LSPs necesitan más información sobre los enlaces que la proporcionada por los protocolos estándar intra-dominio. Estos atributos TE son distribuidos utilizando los mecanismos ya existentes en los protocolos IGP [29], [34].

GMPLS extiende dos protocolos «link-state» intra-dominio tradicionales OSPF-TE e IS-IS-TE. Esto es necesario para codificar y transportar uniformemente la información de un enlace TE. Un enlace TE, por definición, es una representación en los anuncios de estados de los enlaces de los protocolos OSPF e IS-IS y en la base de datos de estados de los enlaces de ciertos recursos físicos y sus propiedades entre dos nodos GMPLS.

GMPLS amplía los protocolos de señalización RSVP-TE y CR-LDP aunque no especifica cuál de ellos debe ser utilizado. La nueva señalización deberá soportar la creación de rutas específicas (*source routing*), transportar los parámetros requeridos del LSP (ancho de

banda, tipo de señal, protección y/o restauración deseada, posición en un determinado multiplex, etc.) para los nuevos tipos de interfaces.

GMPLS introduce mejoras en la escalabilidad como la creación de LSPs jerárquicos, la agrupación de enlaces y los enlaces no numerados [33], [34].

1.5.5 Mejoras de escalabilidad en GMPLS.

En las capas TDM, LSC y FSC es posible tener varios cientos de enlaces físicos paralelos que conectan dos nodos. Esto introduce nuevas restricciones en los modelos de direccionamiento y enrutado IP.

Los nuevos sistemas DWDM permitirán varios cientos de longitudes de onda por fibra.

Se hace impracticable asociar una dirección IP a cada extremo de cada enlace físico para representar cada enlace como una adyacencia de enrutamiento distinta y para anunciar y mantener estados de enlaces para cada uno de estos enlaces. Con este propósito GMPLS amplía los modelos de enrutado y direccionamiento, para aumentar su escalabilidad.

Se pueden utilizar dos mecanismos para aumentar la escalabilidad del direccionamiento y enrutado: enlaces no numerados y agrupamiento de enlaces. Estos mecanismos se pueden combinar. Para implementarlos se necesitan extensiones en los protocolos de señalización (RSVP-TE y CR-LDP) y enrutamiento (OSPF-TE e IS-IS-TE) [15], [33], [35].

1.6 Conclusiones parciales del capítulo.

En este capítulo se expusieron las características de las principales tecnologías de transporte que han sido consideradas a lo largo del desarrollo de las redes de Próxima Generación. Se hizo énfasis en las más utilizadas por las redes de telecomunicaciones a gran escala.

.

CAPÍTULO II: ELEMENTOS DEL NÚCLEO DE LA RED

2.1 Componentes de MPLS.

Cuando se diseña una red con arquitectura MPLS se debe tener en cuenta que como sistema comprende tres conceptos fundamentales: (LSR), (LSP,) y LP, En su forma más simple un sistema MPLS es un conjunto de LSR que envían paquetes etiquetados a través de los LSP.

Los LSR y los LER, son los nodos principales que componen una red MPLS. Los dos son físicamente el mismo dispositivo, un *router* o un *switch* de red troncal que incorpora el *software* MPLS, siendo el administrador el que lo configura para cualquiera de los dos modos de trabajo.

A continuación, haciendo referencia a la Figura 1.3, se muestran los componentes que integran la estructura de un núcleo de red MPLS.

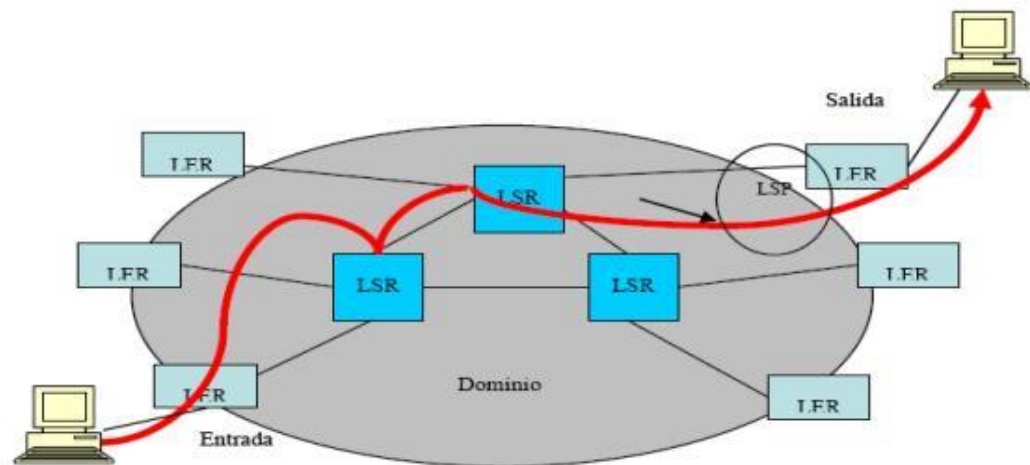


Figura 1.3 Componentes de una red MPLS.

2.1.1 LER.

Los LER se encuentran ubicados en el borde de la red MPLS y desempeñan las funciones de encaminamiento tanto para un dominio MPLS como para un dominio no MPLS (otras redes). El propósito de los LER es el análisis y clasificación del paquete IP que entra a la red de acuerdo a criterios, a esta clasificación por conjuntos de paquetes se le denomina FEC. Una vez analizado el paquete IP se añade una cabecera MPLS y en uno de sus campos denominado Etiqueta se le asigna un valor de acuerdo a su clasificación FEC.

Al salir del dominio MPLS el LER de salida es el que direcciona el paquete a la red de destino por enrutamiento convencional eliminando la cabecera MPLS. El LER de ingreso a la red o dominio MPLS también se lo conoce como *Ingress* LSR y el LER de salida se le llama *Egress* LSR [4], [7], [36].

2.1.2 LSR.

El LSR se encuentra ubicado en el núcleo de la red MPLS, realiza encaminamiento basándose en la conmutación de etiquetas. Una vez que le llega un paquete a una de sus interfaces éste lee la etiqueta de entrada en la cabecera MPLS y busca en la tabla de conmutación la etiqueta y la interfaz de salida para designar la nueva etiqueta que indica el siguiente salto dentro del dominio y finalmente reenvía el paquete por el camino ya designado en el LER (según el FEC).

La conmutación es muy rápida ya que los LSR solo se encargan de la lectura e intercambio de etiquetas obviando la lectura de las cabeceras IP de los paquetes pero es posible que los LSR sean los que retiran la cabecera MPLS en el penúltimo salto antes de salir el paquete por un LER, este hecho puede suceder cuando en un dominio MPLS hay mucho tráfico y resulta mayor procesamiento para el LER, este mecanismo se denomina “remoción en el penúltimo salto” su siglas en inglés PHP.[4], [11], [36].

2.1.3 LSP.

El LSP es una ruta de tráfico específica a través de la red MPLS que sigue un grupo de paquetes que pertenecen al mismo FEC. Esta ruta se crea concatenando los saltos que dan los paquetes para el intercambio de etiquetas en los LSR y para esto utiliza mensajes LDP.

Las rutas LSP se forman desde el destino hacia el origen debido a que el LSR de origen genera las peticiones para crear un nuevo LSP mientras que el destino responde a estas solicitudes formándose de esta manera el LSP hasta el origen. Existen dos métodos para el establecimiento de los LSPs:

- Ruta explícita.
- Salto a Salto.

En la Ruta explícita a partir del primer LSR de salto se construye una lista de saltos específica utilizando los protocolos de señalización o de distribución de etiquetas (RSVP, LDP)

Mientras que en el Salto a Salto cada LSR selecciona el próximo salto según el FEC que esté disponible. El encaminamiento del LSP se realiza mediante protocolos de enrutamiento que utilizan algoritmos de estado de enlace para conocer la ruta trazada completa y tener rutas alternativas si algún enlace falla [35], [37], [38].

2.1.4 FEC.

El FEC es un conjunto de paquetes que son reenviados sobre un mismo camino a través de la red LSP y se determina una vez a la entrada a la red MPLS en un router LER. Para clasificar a los paquetes dentro de un mismo FEC se hace en base a criterios como:

- Dirección IP de origen, destino o direcciones IP de la red.
- Número de puerto de origen o destino
- Campo protocolo de IP, TCP, UDP, ICMP, etc.)
- Valor del campo DSCP de Diff-Serv
- Etiqueta de flujo en IPv6

Cada FEC tiene QoS debido a que se debe tratar a los paquetes que van por el mismo camino de diferente manera, dando prioridad según la necesidad, de manera que se utilizan los recursos de la red óptimamente [17], [39], [40].

2.1.4.1 Agregación.

La Agregación es un mecanismo que permite agrupar varios FEC mediante la asignación de una sola etiqueta para todos, de esta manera se reduce el tiempo de envío de los FEC porque se elimina asociaciones etiqueta/FEC redundantes.

Puede ser posible la agregación cuando a un LSR le llegan desde un mismo LER varios FEC con el mismo origen y destino dentro de la red MPLS asignados al mismo camino LSP [20], [22].

2.1.5 LDP.

El LDP define los mecanismos para la distribución de etiquetas, permite a los LSR descubrirse e intercambiar información sobre las asociaciones FEC/Etiqueta que se han realizado y sobre todo para mantener la coherencia de las etiquetas utilizadas para los distintos tipos de tráfico que conmutan. Con este protocolo se evita que a un LSR le llegue tráfico con una etiqueta que no se encuentra en su tabla, con esto se asegura la rapidez en la conmutación de los LSR.

Para establecer la ruta LSP los LER/LSR establecen sesiones a través de mensajes en los cuales se solicita:

- A su vecino que le informe sobre que etiqueta debe usar para el envío del tráfico por una determinada interfaz, es decir que la distribución de etiquetas se realiza contraria al camino que sigue el tráfico.
- Un LER/LSR informa de las asociaciones Etiqueta/FEC a sus vecinos que las almacenan en sus tablas sin haber solicitado la información, este mecanismo es más eficaz ya que así todos los vecinos LER/LSR mantienen las tablas actualizadas (del mismo LSP) y haciendo el proceso de conmutación de etiquetas mucho más rápido pero incrementando el tráfico de control [16], [38].

2.1.6 LIB.

Un LSR o LER tiene dos tablas, una dedicada a la información de enrutamiento y la segunda con la información a nivel local de las etiquetas conocida como LIB. Los datos de la tabla LIB se relacionan con las etiquetas que han sido asignadas por un LER/LSR y de las asociaciones etiqueta/FEC recibidas de los vecinos del dominio MPLS mediante los protocolos de Distribución de Etiquetas. La información que proporciona una tabla LIB da a conocer sobre la interfaz y etiqueta de entrada seguida de la interfaz y el valor de etiqueta de salida, este proceso se realiza en cada salto de un LSR o LER y permite mantener actualizadas las rutas LSP [20].

2.2 Funcionamiento del envío de paquetes en MPLS.

La base del MPLS está en la asignación e intercambio de etiquetas ya expuesto, que permiten el establecimiento de los caminos LSP por la red. Los LSPs son simplex por naturaleza (se establecen para un sentido del tráfico en cada punto de entrada a la red); el tráfico dúplex requiere dos LSPs, uno en cada sentido. Cada LSP se crea a base de concatenar uno o más saltos en los que se intercambian las etiquetas, de modo que cada paquete se envía de un "conmutador de etiquetas" LSR a otro, a través del dominio MPLS. Un LSR es un router especializado en el envío de paquetes etiquetados por MPLS.

Al igual que en las soluciones de conmutación multinivel, MPLS separa las dos componentes funcionales de control (*routing*) y de envío (*forwarding*). Del mismo modo, el envío se implementa mediante el intercambio de etiquetas en los LSPs. Sin embargo, MPLS no utiliza ninguno de los protocolos de señalización ni de encaminamiento definidos por el "ATM Fórum"; en lugar de ello, en MPLS o bien se utiliza el protocolo RSVP o bien el estándar de señalización LDP. De acuerdo con los requisitos del IETF, el transporte de datos puede ser cualquiera. Si éste fuera ATM, una red IP habilitada para MPLS es ahora mucho más sencilla de gestionar que la solución clásica IP/ATM. Ahora ya no hay que administrar dos arquitecturas diferentes a base de transformar las direcciones IP y las tablas de encaminamiento en las direcciones y el encaminamiento ATM. Esto lo resuelve el procedimiento de intercambio de etiquetas MPLS. El papel de ATM queda restringido al mero transporte de datos a base de celdas. Para MPLS esto es indiferente, ya que puede utilizar otros transportes como Frame Relay, o directamente sobre líneas punto a punto [4], [11], [35].

Un camino LSP es el circuito virtual que siguen por la red todos los paquetes asignados a la misma FEC, al primer LSR que interviene en un LSP se le denomina de entrada o de cabecera y al último se le denomina de salida o de cola. Los dos están en el exterior del dominio MPLS. El resto, entre ambos, son LSRs interiores del dominio MPLS. Un LSR es como un router que funciona a base de intercambiar etiquetas según una tabla de envío. Esta tabla se construye a partir de la información de encaminamiento que proporciona la componente de control. Cada entrada de la tabla contiene un par de etiquetas entrada/salida correspondientes a cada interfaz de entrada, que se utilizan para acompañar a cada paquete

que llega por ese interfaz y con la misma etiqueta (en los LSR exteriores **sólo** hay una etiqueta, de salida en el de cabecera y de entrada en el de cola). En la figura 2.2 se ilustra un ejemplo del funcionamiento de un LSR del núcleo MPLS. A un paquete que llega al LSR por el interfaz 3 de entrada con la etiqueta 45 el LSR le asigna la etiqueta 22 y lo envía por el interfaz 4 de salida al siguiente LSR, de acuerdo con la información de la tabla.

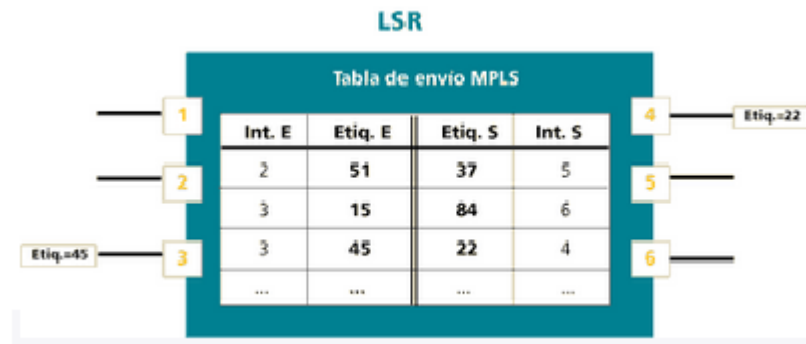


Fig. 2.2: Tabla de envío de un LSR [4].

El algoritmo de intercambio de etiquetas requiere la clasificación de los paquetes a la entrada del dominio MPLS para poder hacer la asignación por el LSR de cabecera. En la figura 2.3 el LSR de entrada recibe un paquete sin etiquetar cuya dirección de destino es 212.95.193.1. El LSR consulta la tabla de encaminamiento y asigna el paquete a la clase FEC definida por el grupo 212.95/16. Asimismo, este LSR le asigna una etiqueta (con valor 5 en el ejemplo) y envía el paquete al siguiente LSR del LSP. Dentro del dominio MPLS los LSR ignoran la cabecera IP; solamente analizan la etiqueta de entrada, consultan la tabla correspondiente (tabla de conmutación de etiquetas) y la reemplazan por otra nueva, de acuerdo con el algoritmo de intercambio de etiquetas. Al llegar el paquete al LSR de cola (salida), ve que el siguiente salto lo saca de la red MPLS; al consultar ahora la tabla de conmutación de etiquetas quita **ésta** y envía el paquete por **ruteo** convencional.

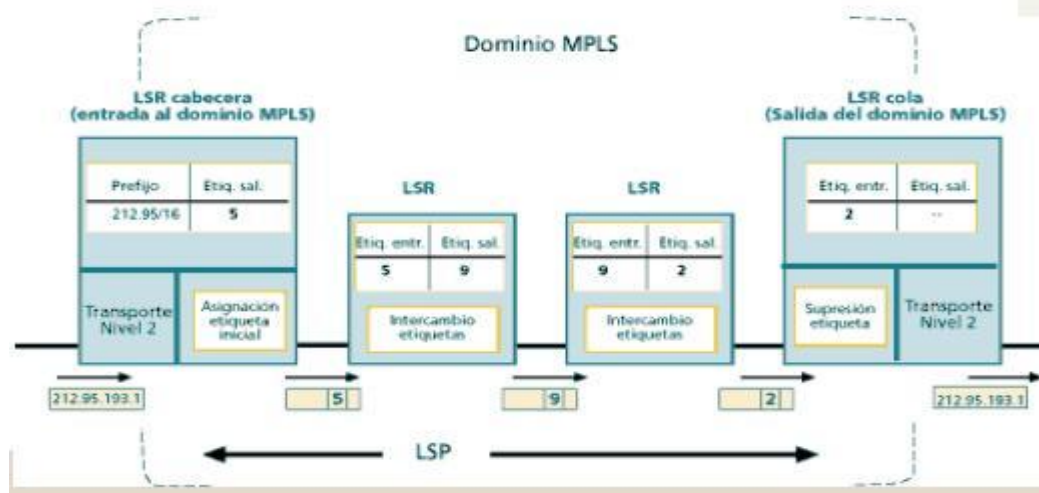


Fig.2.3: LSR de entrada recibe un paquete normal (sin etiquetar).

La identidad del paquete original IP queda enmascarada durante el transporte por la red MPLS, que no "mira" sino las etiquetas que necesita para su envío por los diferentes saltos LSR que configuran los caminos LSP. Las etiquetas se insertan en cabeceras MPLS, entre los niveles 2 y 3. Según las especificaciones del IETF, MPLS debía funcionar sobre cualquier tipo de transporte: PPP, LAN, ATM, Frame Relay, etc. Por ello, si el protocolo de transporte de datos contiene ya un campo para etiquetas se utilizan esos campos nativos para las etiquetas. Sin embargo, si la tecnología de nivel 2 empleada no soporta un campo para etiquetas enlaces PPP o LAN, entonces se emplea una cabecera genérica MPLS de 4 octetos, que contiene un campo específico para la etiqueta y que se inserta entre la cabecera del nivel 2 y la del paquete (nivel 3).

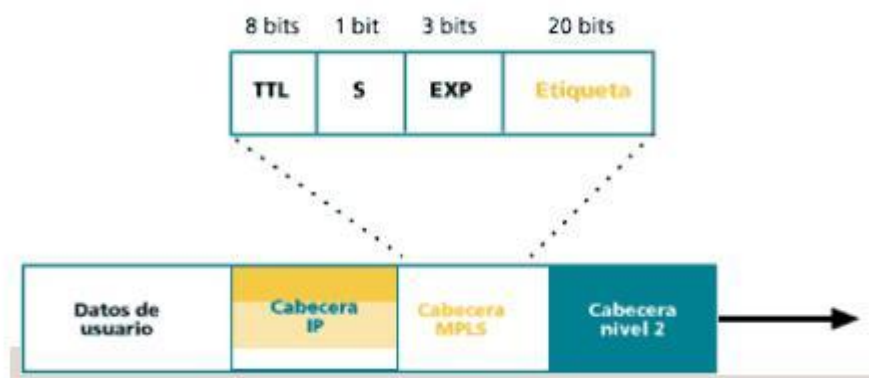


Fig: 2.4 Esquema de los campos de la cabecera genérica MPLS [4].

2.3 Mecanismos de encolamiento.

En los nodos fronteras e interiores de una red se llevan a cabo, entre otras, funciones de encolamiento para acondicionar el tráfico, ya sea cuando entra a la red, durante su recorrido por la misma o simplemente cuando la abandonan. Además, mediante estas funciones de encolamiento los nodos pueden definir el tratamiento particular que recibirán los paquetes de un determinado flujo de tráfico.

Las colas de espera juegan un papel fundamental, ya que mientras mayor tiempo pasen los paquetes en la cola de espera, mayor será el tiempo total de comunicación.

Existen diferentes mecanismos de encolamiento basados en algoritmos que van desde los más simples hasta los sumamente complejos. Cada mecanismo de encolamiento tiene sus ventajas y desventajas, así como escenarios donde es más recomendable usar ese mecanismo en particular, la elección del mecanismo de encolamiento a utilizar depende de lo que se quiera lograr.

2.3.1 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. Este 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 [41], [42].

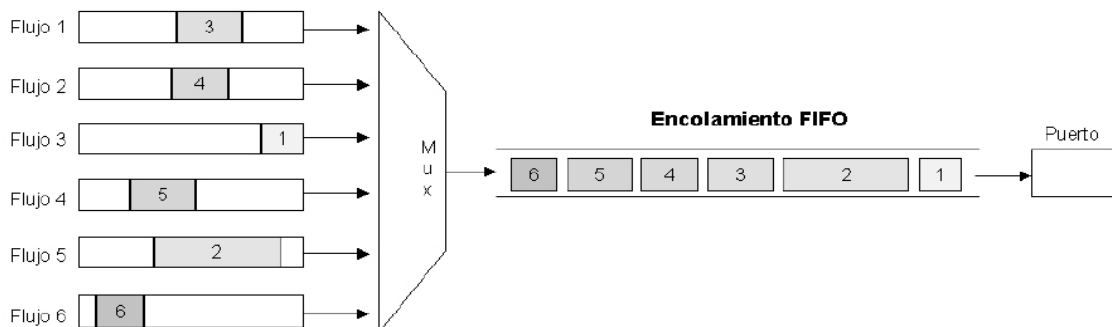


Fig: 2.5 FIFO (First-In First-Out) [41].

2.3.2 WFQ.

Esta disciplina de encolado que es una aproximación del método GPS divide el tráfico entrante en N colas. El ancho de banda del puerto de salida se divide entre las N colas en función de pesos asignados a cada cola. La suma de los pesos siempre tiene que ser 100 %. En WFQ a diferencia de FQ en el cual se recorren las N colas enviando un paquete de cada cola si ésta no está vacía, el planificador saca de las N colas los paquetes en orden de su tiempo de finalización. Este tiempo de finalización se estima haciendo una aproximación al modelo teórico *Weighted bit-by-bit round robin*. Este modelo recorre las colas usando *round robin* sacando cada vez un bit. Cuando un paquete está completo se envía de manera que los paquetes grandes esperan más tiempo a ser enviados que los pequeños. Esta aproximación es sólo teórica y no es aplicable ya que reensamblar los bits de cada trama una vez extraídos de las colas requeriría un procesamiento mayor.

Esta disciplina solventa los problemas que presentaba FQ y permite un reparto diferenciado por clases del ancho de banda de salida. En una situación en la que todas las colas tienen tráfico, el ancho de banda de un flujo queda limitado, lo cual impone una limitación a la hora de medir el ancho de banda en sistemas que tengan este tipo de encolado [43], [44].

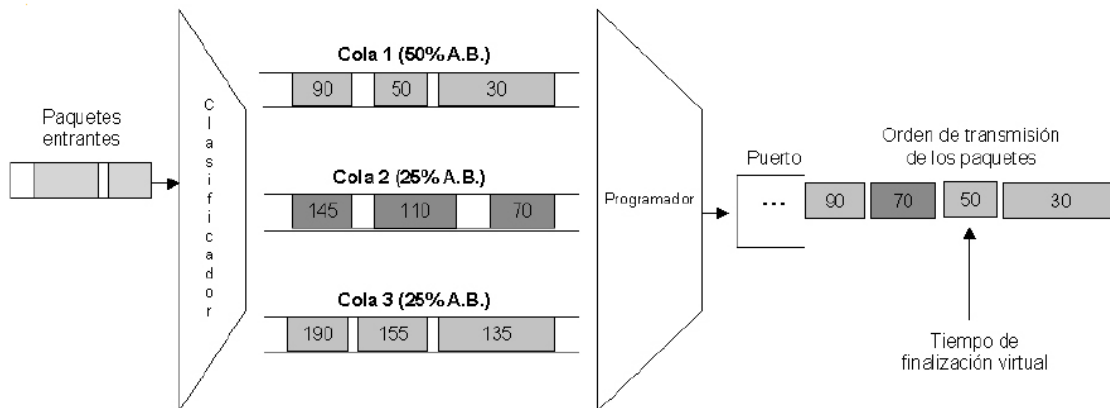


Fig: 2.6 Funcionamiento de WFQ [43].

2.3.3 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 analiza. 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 [45].

2.3.4 PQ.

El Encolamiento de Prioridad 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 de una cola de mayor prioridad ingresa, esta última es atendida inmediatamente. Este mecanismo se ajusta a condiciones donde existe un tráfico importante, pero puede causar la total falta de atención a colas de menor prioridad [46].

PQ ofrece los siguientes beneficios:

- Para routers basados en software, el encolamiento *PQ* supone una carga mucho menor en el sistema en comparación con disciplinas de servicio de colas más elaboradas.
- *PQ* permite a los routers organizar los paquetes almacenados y por tanto servir una clase de tráfico de modo diferente a otras. Por ejemplo, se pueden colocar prioridades a las aplicaciones de tiempo real, como voz y video interactivo, y que se traten de forma prioritaria frente a otras aplicaciones que no operan en tiempo real.

Pero *PQ* también tiene varias limitaciones:

- Si la cantidad de tráfico de alta prioridad no se acondiciona, el tráfico de baja prioridad, puede experimentar un retardo excesivo mientras espera a que se sirva el tráfico de alta prioridad.
- Si el volumen de tráfico de alta prioridad llega a ser excesivo, se puede descartar el tráfico de baja prioridad cuando las memorias reservadas para este tipo de tráfico se desborden.
- Un mal comportamiento de un flujo de tráfico de alta prioridad, puede añadir un aumento significativo del retardo y de la varianza del retardo (*jitter*) experimentado por otros flujos de tráfico de alta prioridad con los que comparte la cola.
- PQ no es la solución a las limitaciones del encolamiento FIFO en donde se favorecían a los flujos UDP sobre los TCP, durante períodos de congestión.

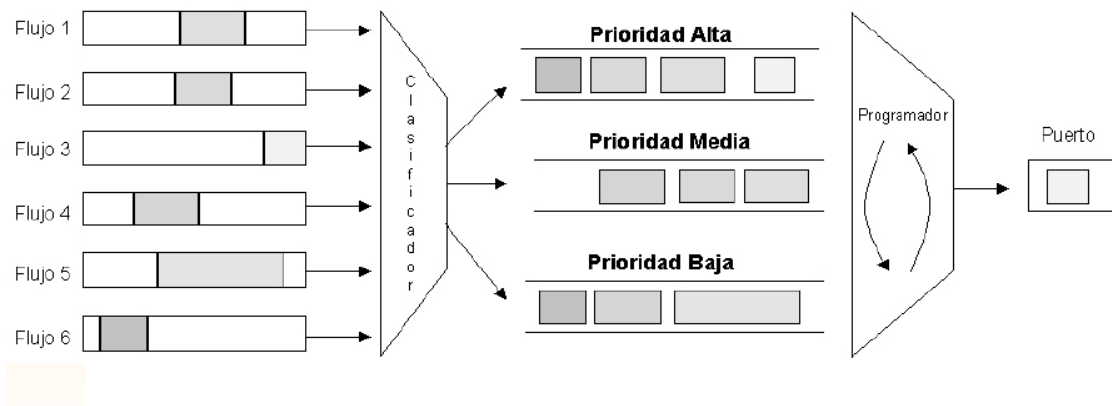


Figura 2.7 Funcionamiento de PQ [46].

2.3.5 CQ.

Para evadir la rigidez de PQ, se opta por utilizar Encolamiento Personalizado CQ. 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. 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 [47].

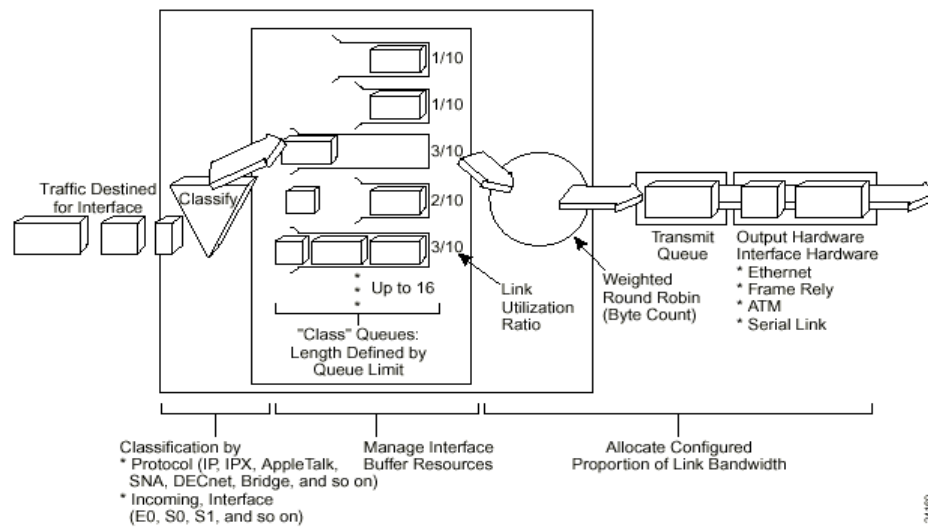


Fig: 2.8 Funcionamiento de CQ .

2.4 Modos de simulación en OPNET.

Dos de los modos de simulación que posee el OPNET Modeler son: DES y análisis de flujo, los cuales presentan características y facilidades totalmente diferentes.

El DES ofrece modelos altamente detallados, que explícitamente simula los mensajes de protocolos y paquetes. Los modelos DES ejecutan el protocolo en la misma forma que en un ambiente de producción. Aunque DES ofrece buenos resultados, las corridas de la simulación demoran más que las corridas del método análisis de flujos. En una simulación DES se utilizan eventos para describir sucesos o acciones que tienen lugar en un determinado momento. Cada uno de estos eventos tiene un instante de incidencia puntual en la escala temporal. Esta escala, al igual que el resto de magnitudes, es discreta. Durante la simulación, se puede distinguir entre el tiempo real y el tiempo de simulación, por ello se utilizan variables contador para representar el momento actual y las cantidades de tiempo. También se utilizan varias variables de estado para representar la fase del sistema simulado. Por lo que se refiere al sistema, evoluciona en la memoria del ordenador, produciéndose diferentes eventos que edifican las variables de estado, y estas a su vez determinan los futuros eventos. Cada evento tiene un instante de incidencia puntual, es decir, pueden ser representados sobre la escala temporal discreta de la simulación ocupando una única posición. De este modo, dichos eventos logran ser ordenados cronológicamente, según su

instante de incidencia, para ser procesados. En este tipo de simulación el evento es la unidad de ejecución. Cada uno describe una acción, y el resultado de esta es la modificación de las variables de estado. Esta característica está especialmente soportada por los lenguajes OOP.

Cuando se ejecuta la simulación de eventos discretos de una red modelada, posteriormente, existen una serie de facilidades que se pueden visualizar a partir de los resultados de las simulaciones. A continuación se mencionan algunas de ellas:

- Permite conocer el retardo, de forma general, que existe en una red.
- Permite conocer el tráfico recibido y enviado, tanto en bytes/s como en *packets/sec*, de diferentes tipos de aplicaciones y protocolos, como por ejemplo: *Email*, FTP, HTTP, Voz,
- Videoconferencia, etc.
- En los nodos de la red, permite conocer el tráfico de una interfaz determinada, así como la representación gráfica de mecanismos de encolamiento de calidad y servicio usados en la red.
- Permite conocer la actividad de convergencia de los protocolos de enrutamiento, las tablas de rutas y la cantidad de saltos entre los nodos seleccionados.
- Brinda la facilidad de conocer la cantidad de conexiones TCP que existen o el tráfico UDP que presenta la red modelada.
- Permite ver estadísticas de enlace como el retardo de colas, el promedio de bits transmitidos o recibidos y la utilización de los enlaces en la red [48].

CAPÍTULO III: PROPUESTA DEL NÚCLEO NGN

En este capítulo se realiza una propuesta de núcleo de red NGN para la provincia de Villa Clara y se analizan los resultados alcanzados, teniendo en cuenta las características de los dispositivos u objetos que brinda la herramienta de modelación y simulación de redes OPNET Modeler 14.0.

3.1 Planteamiento del diseño.

El diseño incluirá la selección de objetos que posee la aplicación OPNET Modeler 14.0 de acuerdo a las características de los ISAM 7302 como parte del backbone, comportándose cada uno de ellos con enrutadores de borde (LER) de la red IP/MPLS. El protocolo IP ofrece varios tipos de servicio según las VPNs creadas. La demanda de servicios garantizados y las nuevas aplicaciones multimedia llevan a la necesidad de establecer grados de QoS. La red MPLS al diseñarse brindará QoS a las diferentes VRF configuradas según los protocolos de enrutamiento seleccionado.

3.2 Diagrama de Topología.

Para la elaboración de del diagrama topológico se tuvo en cuenta la selección de los modelos de nodo *ethernet2_slip8_ler* y *ethernet2_slip8_lsr*, los cuales representan los enrutadores LER y LSR respectivamente. Otra alternativa utilizada fueron los modelos de nodos 100BaseT que simbolizan una red LAN caracterizada internamente por un servidor Ethernet, un conmutador y 250 estaciones. También se presentan los modelos de nodo *ppp_wkstn*, los cuales representan estaciones que implementan tanto servicio de datos como de multimedia. Además de estos dispositivos, se utilizaron enlaces con razones de datos a 148.61 Mbps, dichos enlaces, representan el objeto PPP_SONET_OC3. Otros tipos de configuraciones realizadas en espacio de trabajo del OPNET Modeler fue la creación de 156 LSP dinámicos entre todos los enrutadores ISAM de la red modelada, los cuales caracterizan cada una de las VPN asignadas a estos dispositivos. En el Anexo 14 se visualizan estas configuraciones. La figura 3.1 es un diagrama lógico de la red NGN implementada en Villa Clara.

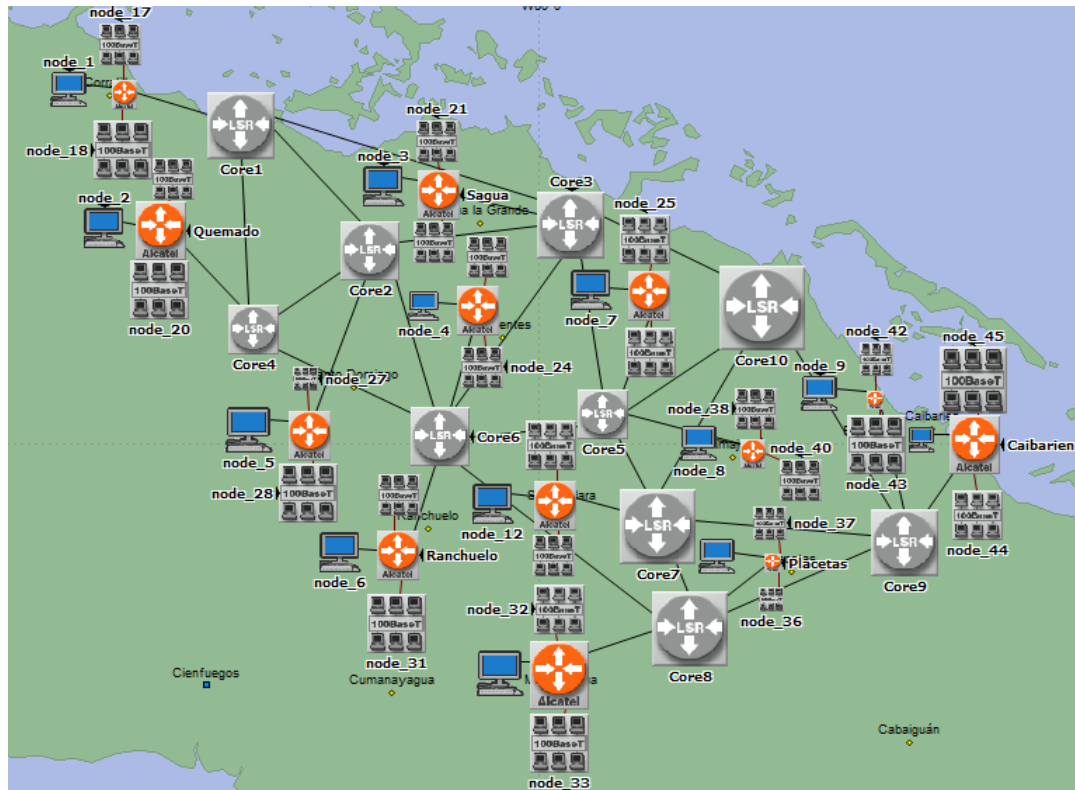


Figura 3.1. Diagrama lógico de la red NGN de Villa Clara.

3.3 Aplicaciones y Perfiles.

Para la creación de los perfiles de la red modelada se tuvieron en cuenta varias aplicaciones generales que representan las principales necesidades de los clientes de una red NGN, entre ellas se encuentran por la parte de audio la definición de aplicación Llamada de voz sobre IP, en multimedia se percibe video conferencia y en la parte de datos aplicaciones imprescindibles tales como *Email*, FTP e implementaciones de sitios *web* utilizando HTTP.

3.4 Configuración de las VPN en la red IP/MPLS.

3.4.1 Creación de las interfaces de lazo cerrado.

Para que los ISAM 7302 soporten el enrutamiento basado en VPN deben incluir un direccionamiento IP en las interfaces de lazo cerrado las cuales asocian los procesos OSPF y BGP, asegurándose de que no se va a perder las sesiones OSPF o BGP por un problema físico en el interfaz, ya que las interfaces de lazo cerrado son interfaces lógicas. También sirven para tener latente el protocolo de ruteo implementado OSPF, el cual al no detectar

interfaces en el dispositivo después de cierto tiempo, lo descarta de la topología de red por este motivo se crean las interfaces de lazo cerrado ya que las mismas nunca se caerán. Por tanto, es necesario establecer en las mismas el protocolo de enrutamiento OSPF para lograr la conmutación de etiquetas en la red MPLS debido a que se implementa el protocolo de distribución de etiquetas (LDP).

3.4.2 Creación de las VRF.

Las instancias de Redes Privadas Virtuales en Reenvío/Enrutamiento se deben crear en todos los ISAM o LER de la red IP/MPLS, las cuales, una vez creadas, serán asignadas a las interfaces específicas de las subredes las cuales representan audio, video y datos. Dichas instancias VRF poseen un valor de ruta que las distingue a cada una debido a que esta cuantía es única para las VPN que componen la red y está configurada para poseer dos valores separados por puntos suspensivos. La tabla 3.1 muestra las tres VRF creadas en todos los dispositivos ISAM 7302 y sus respectivos valores.

Tabla 3.1. Tablas VRF para los enrutadores ISAM 7302.

Nombre de las VRF	Ruta Distintiva
VPN_Audio	1:10
VPN_Video	1:20
VPN_Datos	1:30

3.4.3 Asignación de las VRF.

Una vez creadas las instancias VRF estas se asignan a las interfaces correspondientes, por tanto, cada enrutador tiene conectado tres subredes las cuales caracterizan el tráfico de videoconferencia (1:20), voz (1:10), servicios ftp (1:30), servicio http (1:30) y correo (1:30).

3.4.4 Protocolo de enrutamiento OSPF.

Para la elección del protocolo de enrutamiento en la red modelada se tuvieron en cuenta las siguientes especificaciones:

- Enrutamiento Jerárquico
- Enrutamiento dinámico de enlace estado
- Calculo de la métrica del menor costo
- Seguridad y autenticación

Dichas características las posee el protocolo de enrutamiento OSPF, el cual es un protocolo de enrutamiento jerárquico de pasarela interior, de enrutamiento dinámico IGP, que usa el algoritmo *SmothWall Dijkstra* enlace-estado LSE para calcular la ruta más corta posible, utilizando la métrica de menor costo, además, construye una base de datos enlace-estado LSDB idéntica en todos los enrutadores de la zona. OSPF es probablemente el tipo de protocolo IGP más utilizado en redes grandes. IS-IS, otro protocolo de enrutamiento dinámico de enlace-estado, es más común en grandes proveedores de servicio. Puede operar con seguridad usando MD5 para autenticar a sus puntos antes de realizar nuevas rutas y antes de aceptar avisos de enlace-estado.

3.4.5 Border Gateway Protocol.

En la red NGN implementada para la provincia de Villa Clara cada enrutador de borde utiliza BGP, con el objetivo de estar enviando constantemente información sobre los atributos de conexión para las VPN establecidas. Esta información basada en paquetes estará utilizando la interfaz virtual de lazo cerrado la cual es la designada para realizar la conmutación de etiquetas hacia dentro de la red IP/MPLS con el resto de los enrutadores de borde de la red, o sea, el resto de los ISAM 7302 correspondientes a cada municipios. Para que BGP funcione debidamente debe tener creado una serie de familias de direcciones para IPv4 VRF y otras para las VPNv4 que se utilizan en la red.

3.4.6 Familia de Direcciones para IPv4 VRF.

Cada ISAM posee una familia de direcciones asignada a cada instancia VRF con el objetivo de realizar una redistribución por defecto a las rutas conectadas directamente y para redistribuir el protocolo OSPF dentro de la ruta IP/MPLS.

Familia de Direcciones para VPNv4.

En las configuraciones establecidas se crean las familias de direcciones para la VPNv4 en donde se activa las direcciones IPv4 de lazo cerrado del resto de los ISAM

correspondientes a cada municipio que componen la NGN de Villa Clara. Donde se estará enviando constantemente información sobre los atributos de conexión para la VPN establecida.

3.4.7 Calidad de Servicio

La QoS a implementar es CQ. Esta disciplina de encolado que es una aproximación del método GPS (Generalized Processor Sharing) divide el tráfico entrante en N colas. El ancho de banda del puerto de salida se divide entre las N colas en función de pesos asignados a cada cola. Por tanto, cada ISAM tendrá asignado este tipo de Calidad de Servicio basada en clases asignando un porcentaje de ancho de banda en cada enlace priorizando un 75 % a las aplicaciones que adicionan tráfico de voz y video en la red IP/MPLS basada en VPN.

3.5 Evaluación de los resultados de la propuesta.

La figura 3.2 muestra el tráfico, en paquetes por segundo, total enviado por OSPF en todas las interfaces de la red que lo soportan, tanto para las interfaces reales como para las virtuales de lazo cerrado. El pico que comienza en los 105 segundos es debido a que los perfiles configurados en la red modelada se inician bajo la variable “uniform (100, 110)”.



Figura 3.2. Tráfico total OSPF.

La **Figura 3.3** muestra el tráfico total enviado por los paquetes *hello* en paquetes/seg, por todas las interfaces de OSPF. Cada enrutador envía periódicamente a sus vecinos un paquete que contiene el listado de vecinos reconocidos por el enrutador, indicando el tipo de relación que mantiene con cada uno.

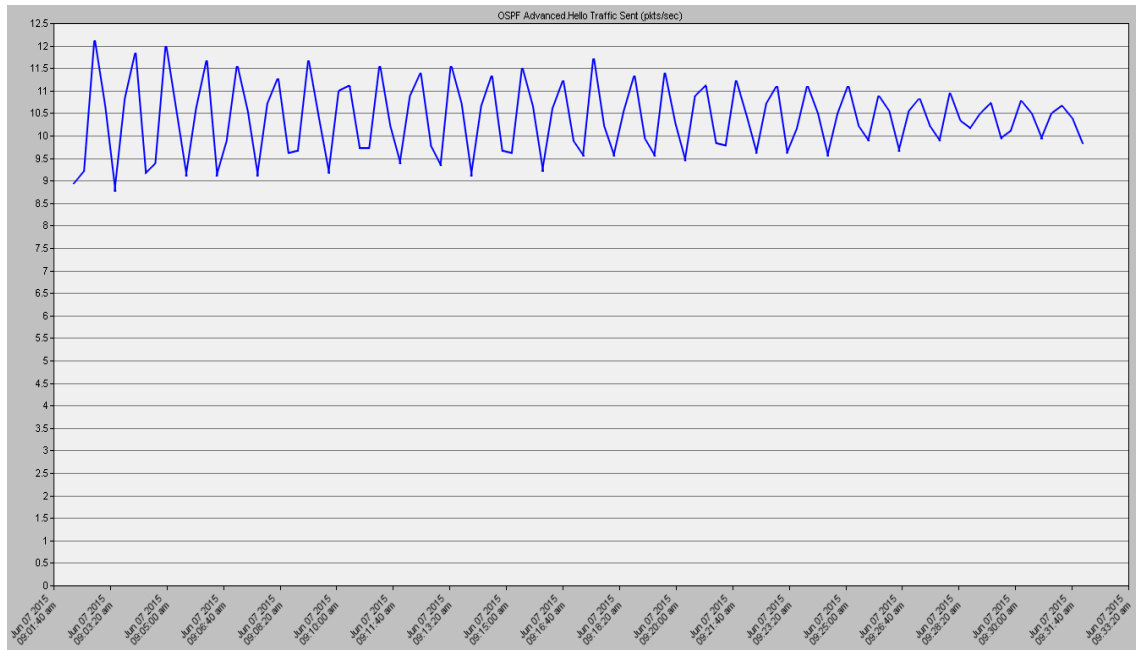


Figura 3.3. Tráfico total de paquetes hello de OSPF.

La **Figura 3.4** muestra el retardo de la QoS CQ basada en el ToS. Esta imagen modela la demora en segundos de los paquetes en cada interfaz del ISAM de Santa Clara utilizando CQ. Esta demora depende del tamaño de cola individual, o sea, del flujo de tráfico y de la razón de datos de los enlaces del enrutador.

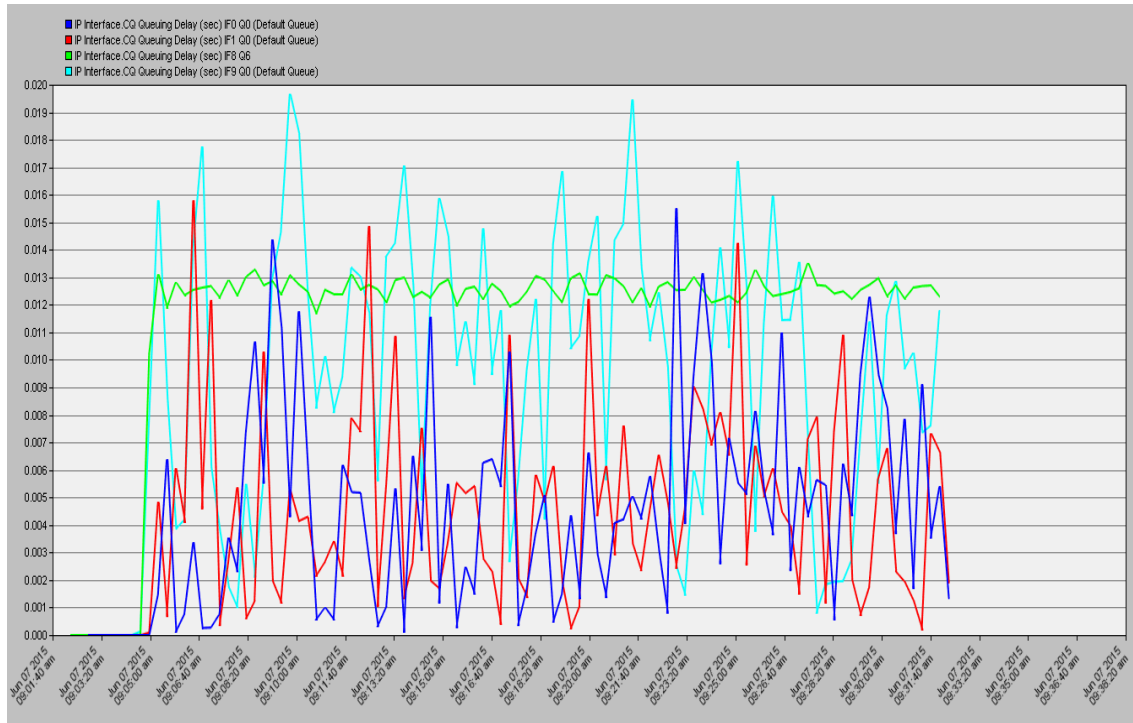


Figura 3.4. Retardo de CQ basado en ToS.

La **Figura 3.5** muestra la variación del retardo de la QoS CQ basada en el tipo de Servicio. Esta imagen modela en segundos la desviación típica de disponer en una cola de espera la demora experimentada por los paquetes de cada interfaz.

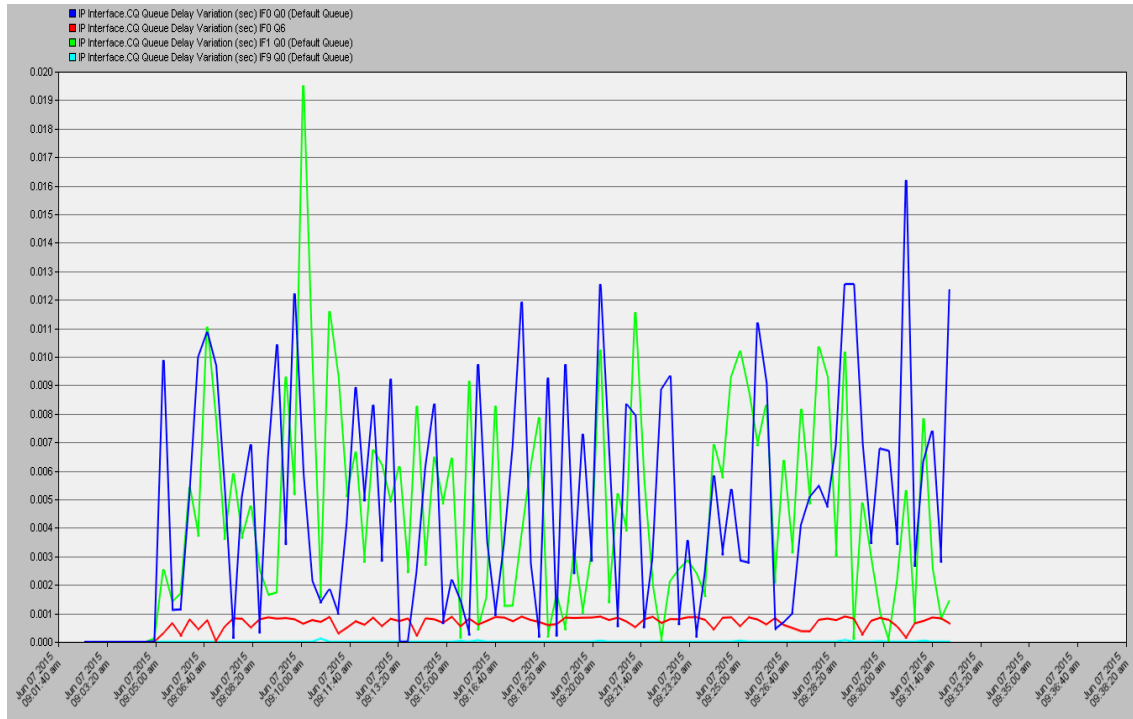


Figura 3.5. Variación del retardo de CQ basado en ToS.

La figura 3.6 muestra la razón de transferencia para todas las interfaces del ISAM 7302 correspondiente al municipio de Santa Clara, donde la suma total se corresponde al flujo de tráfico que existe entre este nodo y un enrutador correspondiente a la núcleo IP/MPLS. Dicha grafica representa el número promedio de paquetes con éxitos recibidos o transmitidos por cada canal y simula el desempeño de las VPN para este ISAM Alcatel.

Lo mismo sucede en la figura 3.7, la cual muestra el tráfico que existe entre los enrutadores LER o ISAM de la red NGN de la Provincia de Villa Clara con los núcleos correspondientes a la red IP/MPLS, añadiendo la excepción del ISAM de Santa Clara, el cual muestra un nivel superior de razón de transferencia debido a que en sus subredes se implementan los servicios de voz, de clientes ftp, de clientes http, email y videoconferencia, utilizando los servidores modelados.

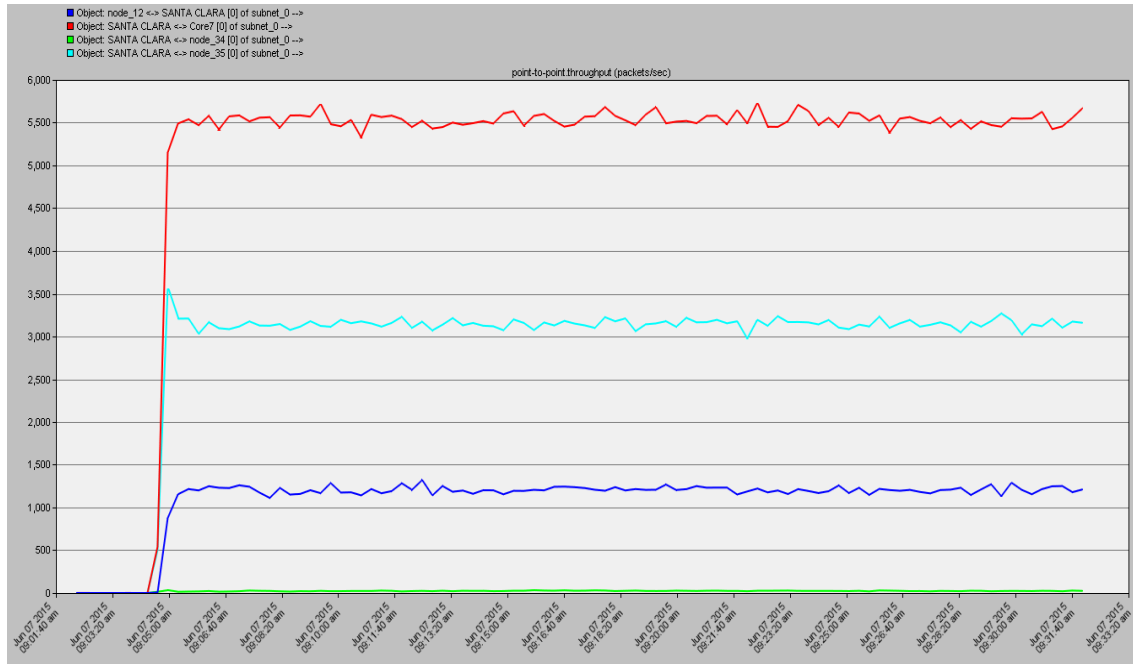


Figura 3.6. Razón de transferencia para las Interfaces del enrutador ISAM correspondiente al municipio de Villa Clara.

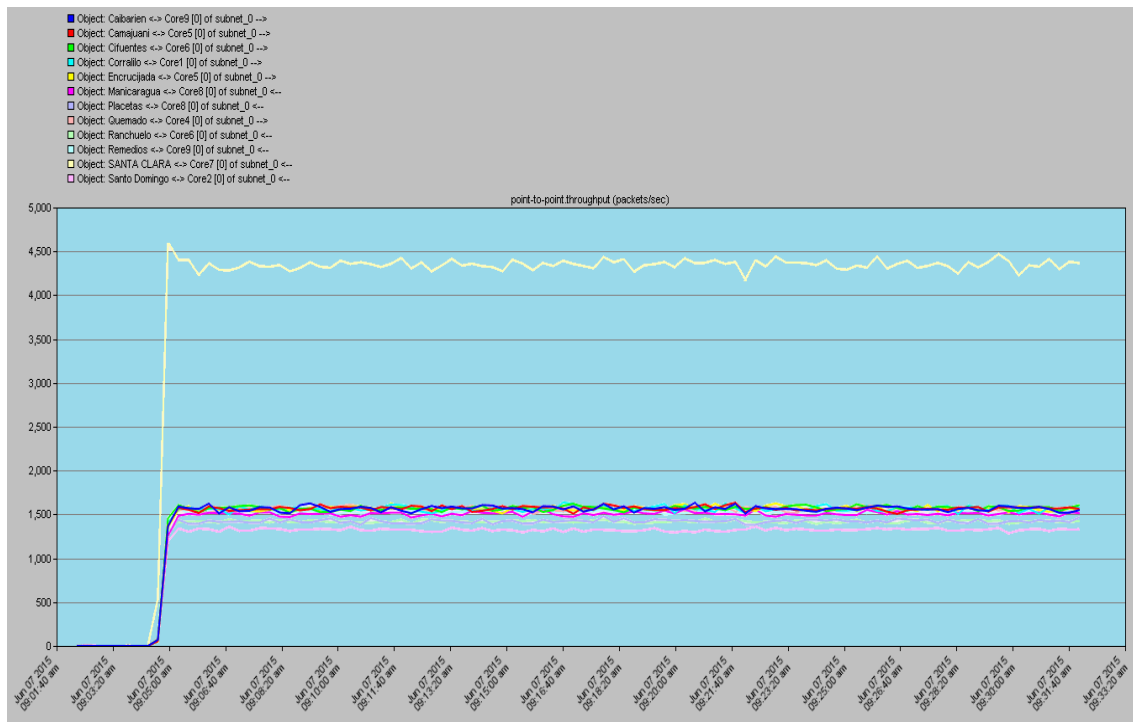


Figura 3.7. Razón de transferencia de los enlaces entre los ISAM y los enrutadores del núcleo de la red IP/MPLS.

La figura 3.8 muestra la carga en bits/seg, de las VPNs para los perfiles de voz, videoconferencia y datos en el enrutador ISAM de Santa Clara. Esta estadística mide la cantidad del tráfico VPN presentado por las interfaces de este dispositivo. Estos datos de informe de estadística solamente se corresponden para las VPN de MPLS que soportan BGP en los enrutadores de borde ISAM 7302.

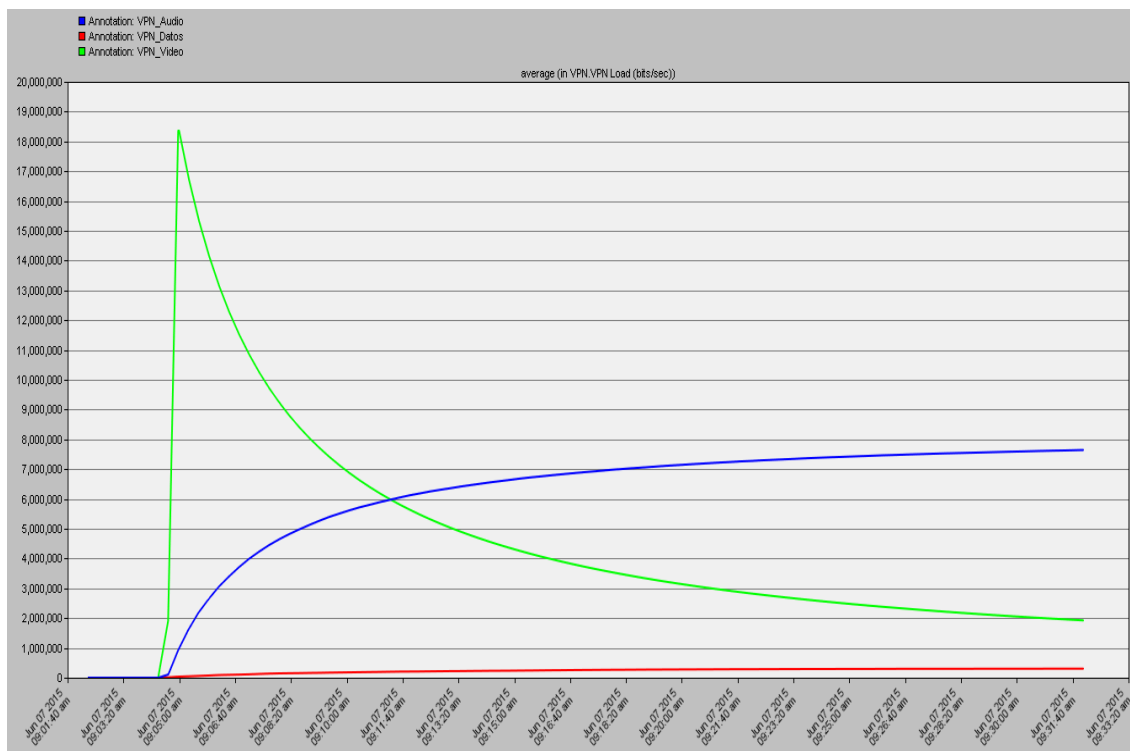


Figura 3.8. Carga de las VPN configuradas en el enrutador ISAM de Santa Clara.

La figura 3.9 muestra el retardo de paquetes para cada LSP correspondiente al enrutador de Santa Clara, en esta estadística se modela la demora experimentada por cada paquete en los LSP de ida entre los enrutadores, seleccionando al ISAM de Santa Clara como fuente.

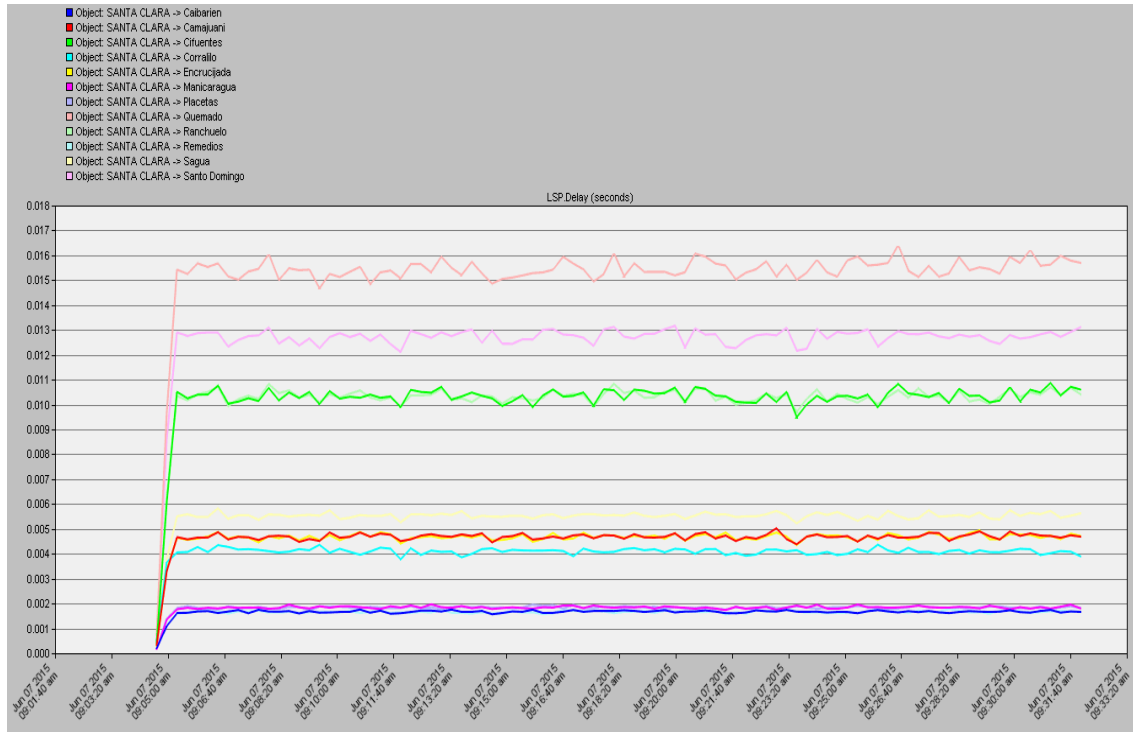


Figura 3.9. Retardo de los LSP conectados al ISAM Santa Clara.

La figura 3.10 muestra el retardo de flujo de tráfico para las VPN en el LSP de Santa Clara a Caibarién. Esta estadística muestra la demora experimentada por cada paquete que pertenece a una circulación específica en el LSP. Además, en esta modelación se indica el tiempo de cada paquete para una conexión en particular dentro de la conmutación de etiquetas en cada trayectoria. Esta estadística está comentada con la dirección IP de la fuente y destino. Dichas direcciones se refieren al retardo de flujo entre las tres conexiones que posee cada dispositivo de enrutamiento en la Red NGN de Villa Clara.

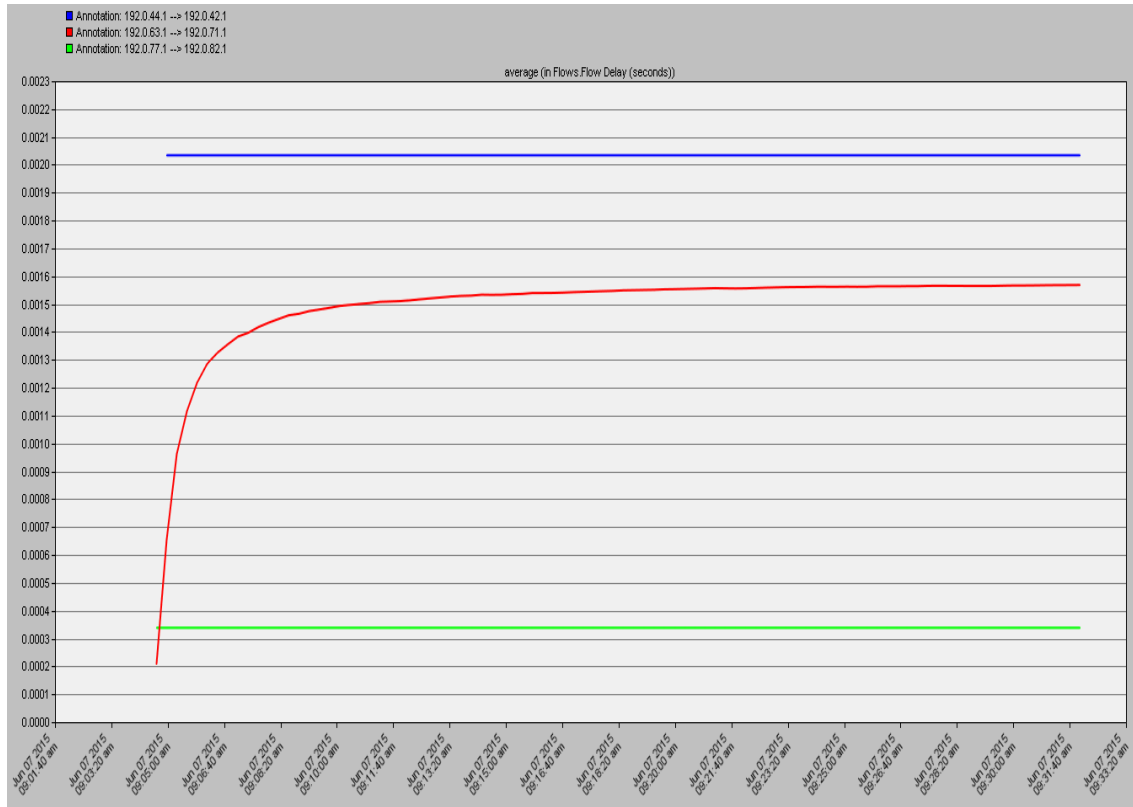


Figura 3.10. Retardo de flujo para las conexiones VPN entre los enrutadores de Caibarién y Santa Clara.

La figura 3.11 muestra el *jitter* para las aplicaciones de Voz configuradas en la red modelada utilizando el tipo de servicio EF. Esta estadística consiste en que si dos paquetes consecutivos dejan el nodo de fuente con tiempos t_1 & t_2 y llegan al nodo destino con tiempos t_3 & t_4 , entonces el *jitter* queda expresado como:

$$\text{Jitter} = (t_4 - t_3) - (t_2 - t_1) [\text{seg}]$$

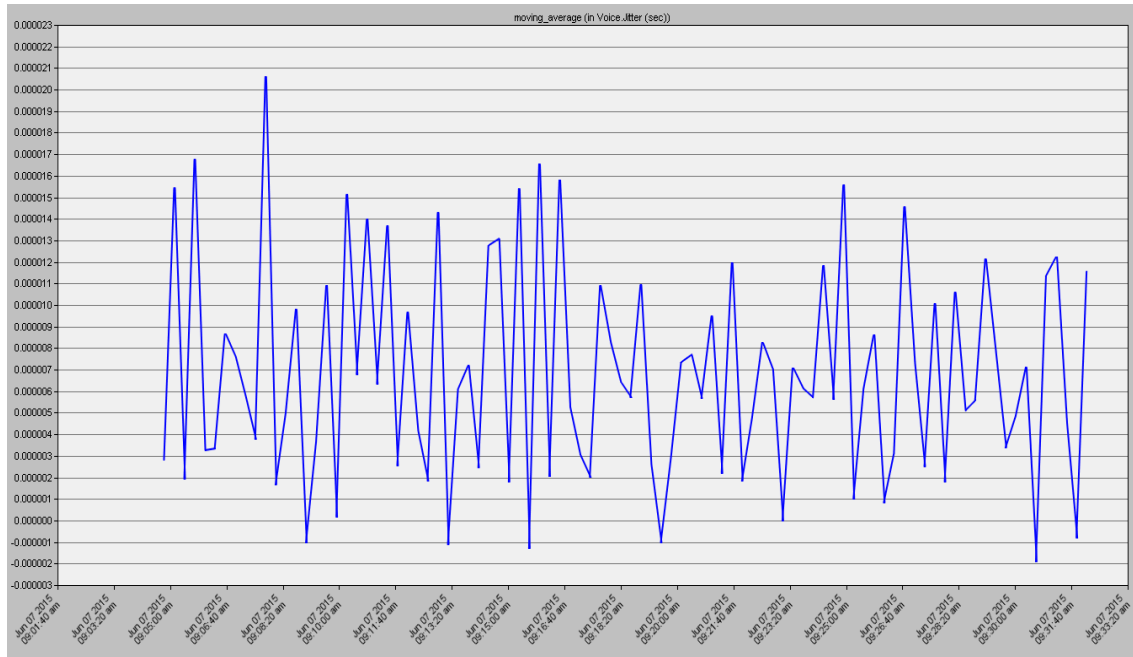


Figura 3.11. *Jitter* de la aplicación de voz_ip_pcm_quality.

La figura 3.12 muestra la variación del retardo de paquetes para la aplicación de voz creada en la red. Dicha aplicación utiliza el estándar de codificación de audio G.711. Esta estadística representa la demora punto a punto para un paquete de voz medido en el momento en que es creado, hasta el momento en que es recibido.

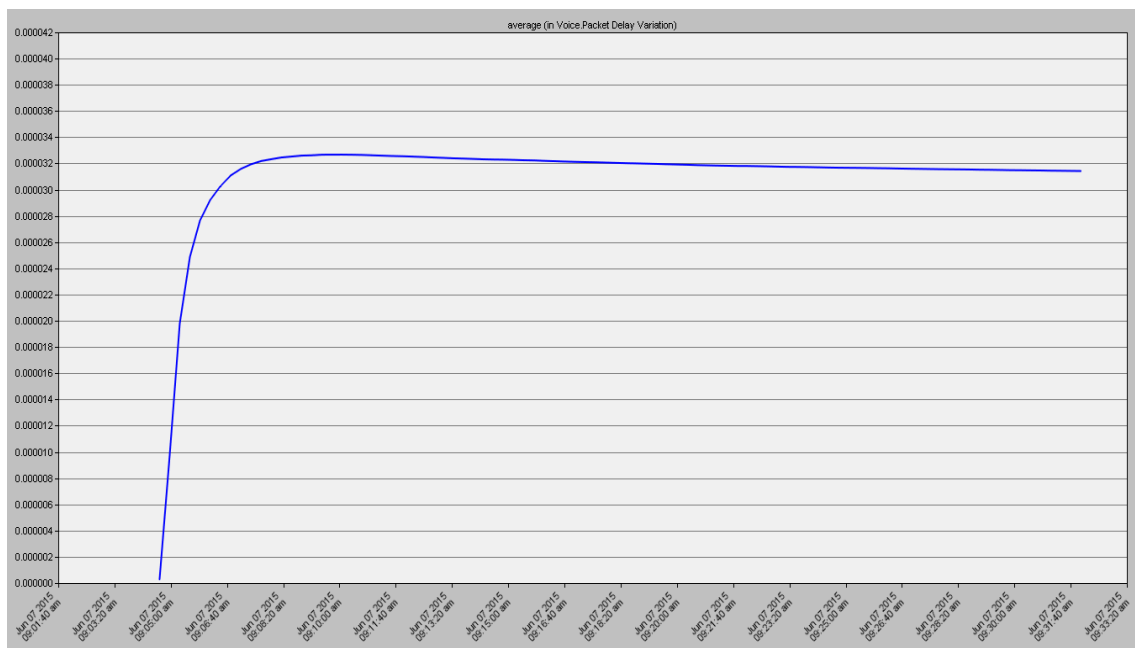


Figura 3.12. Variación del retardo de paquetes de voz.

La figura 3.13 representa el retardo total punto a punto para todos los paquetes recibidos en todas las estaciones de la red. Dicha gráfica es expresada en segundos.

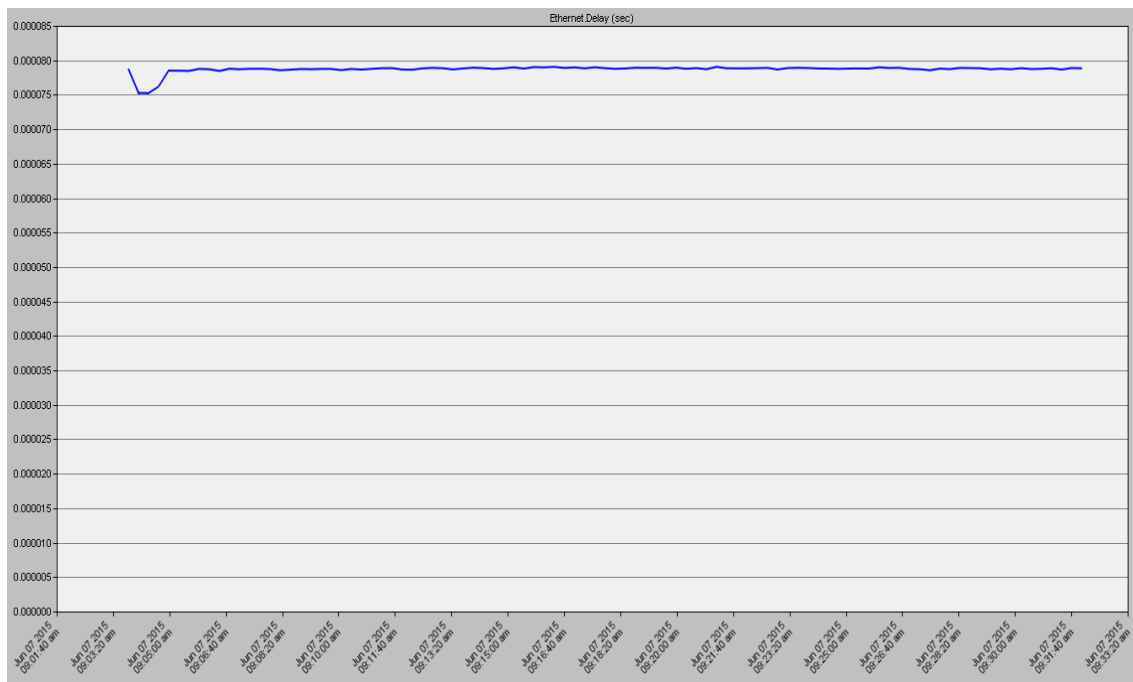


Figura 3.13. Retardo de red para toda la red NGN de Villa Clara.

En los Anexos se muestra el resto de las configuraciones realizadas para la red modelada.

CONCLUSIONES Y RECOMENDACIONES.

Conclusiones:

1. El protocolo MPLS es una tecnología que se desempeña **el núcleo** de una red de telecomunicaciones y permite la aplicación de técnicas de ingeniería de tráfico, la diferenciación de los servicios por clases y garantiza la escalabilidad de la red mediante técnicas de redes privadas virtuales.
2. Una red MPLS está compuesta por enrutadores de borde que se encargan de colocar las etiquetas en los paquetes, los conmutadores de etiquetas (LSR), los caminos virtuales de etiquetas y distintas tecnologías asociadas que garantizan calidad de servicio y un funcionamiento estable de la red.
3. En el núcleo de la red el comportamiento de los parámetros de QoS se mantuvieron en niveles adecuados, la carga de tráfico en la red a partir de los servicios y las técnicas configuradas tuvieron un comportamiento estable, confirmando el alto desempeño del núcleo propuesto.
4. El núcleo de la red NGN de Villa Clara se basa en una solución IP/MPLS con técnicas de QoS, VPN e Ingeniería de Tráfico que garantizan el correcto funcionamiento de los servicios críticos, la escalabilidad y la estabilidad de la red.

REFERENCIAS BIBLIOGRÁFICAS.

- [1] Sector de Normalización de las Telecomunicaciones de la UIT, Serie Y: infraestructura mundial de la información, aspectos del protocolo internet y redes de la próxima generación. Redes de la próxima generación – marcos y modelos arquitecturales funcionales. 2004.
- [2] A. M. Restrepo Restrepo, «Visión general de las redes de próxima generación (NGN)», Trabajo de Grado, Universidad Pontificia Bolivariana, Venezuela, 2013.
- [3] M. C. C. Reyes, J. D. D. M. Sánchez, y A. F. F. C. Gómez, «Aspectos Técnicos y Proceso de Interconexión a Nivel de Servicios en Las NGN», REVISTA ENTRE CIENCIA E INGENIERÍA, n.º 8, pp. 21–36, 2011.
- [4] S. K. Narváez Pupiales, «Diseño de una red de backbone con tecnología MPLS para el soporte de servicios triple play en la empresa Ecuonet-Megadatos SA», 2011.
- [5] J. O. Ordoñez Ordoñez y J. L. Inga Lojano, «Análisis de ubicación, y manual para montaje e instalación de nodos activos en el sector urbano del cantón Cuenca, para la Empresa de Telecomunicaciones, Agua Potable, Alcantarillado y Saneamiento ETAPA EP», 2014.
- [6] G. García Girón, «Propuesta de migración de la red NGN de una operadora implementada en IP hacia MPLS», 2011.
- [7] M. O. Tapasco García, «MPLS, el presente de las redes IP», 2008.
- [8] D. O. Awduche, «MPLS and traffic engineering in IP networks», Communications Magazine, IEEE, vol. 37, n.º 12, pp. 42–47, 1999.
- [9] E. Rosen, A. Viswanathan, y R. Callon, «RFC 3031 Multiprotocol Label Switching Architecture,(January 2001)», Status: STANDARDS TRACK.
- [10] S. Sánchez López, Interconnection of IP/MPLS Networks Through ATM and Optical Backbones using PNNI Protocols. Universitat Politècnica de Catalunya, 2003.
- [11] A. Delfino, S. Rivero, y M. San Martín, «Ingeniería de tráfico en Redes MPLS», 2006.
- [12] M. K. Girish, B. Zhou, y J.-Q. Hu, «Formulation of the traffic engineering problems in MPLS based IP networks», 2000, pp. 214–219.
- [13] N. Rouhana y E. Horlait, «Differentiated services and integrated services use of MPLS», 2000, pp. 194–199.
- [14] G. L. Perafán y V. S. Muñoz, «Hacia la NGN en Colombia», Gerencia Tecnologica Informatica, vol. 10, n.º 26, 2011.
- [15] S. A. Thomas, «IP Switching and Routing Essentials: Understanding RIP, OSPF, BGP, MPLS, CR-LDP and RSVP-TC», Wiley. New York, pp. 263–264, 2001.
- [16] G. Ahn y W. Chun, «Design and implementation of MPLS network simulator supporting LDP and CR-LDP», 2000, pp. 441–446.

- [17] I. F. Akyildiz, T. Anjali, L. Chen, J. C. de Oliveira, C. Scoglio, A. Sciuto, J. A. Smith, y G. Uhl, «A new traffic engineering manager for DiffServ/MPLS networks: design and implementation on an IP QoS Testbed», *Computer Communications*, vol. 26, n.º 4, pp. 388–403, 2003.
- [18] J.-M. Chung, «Analysis of MPLS traffic engineering», 2000, vol. 2, pp. 550–553.
- [19] E. D. Osborne y A. Simha, *Traffic engineering with MPLS*. Cisco Press, 2002.
- [20] S. Yasukawa, «Signaling requirements for point-to-multipoint traffic-engineered MPLS label switched paths (LSPs)», 2006.
- [21] A. Sprintson, M. Yannuzzi, A. Orda, y X. Masip-Bruin, «Reliable routing with QoS guarantees for multi-domain IP/MPLS networks», en *INFOCOM 2007. 26th IEEE International Conference on Computer Communications*. IEEE, 2007, pp. 1820–1828.
- [22] C. Ochoa, J. C. Cuellar, Y. Díaz, A. Navarro, F. G. Guerrero, y D. Blandón, *Medición de la Calidad de Servicio en Redes NGN.*, 1era edición. Colombia: CINTEL, 2010.
- [23] I. Pepelnjak y J. Guichard, *MPLS and VPN architectures*, vol. 1. Cisco Press, 2002.
- [24] A. Barbir, N. Mistry, y W. Ding, *Selection techniques for logical grouping of VPN tunnels*. Google Patents, 2004.
- [25] R. Flores Moyano y S. González Martínez, «Protocolo múltiple por conmutación de etiquetas (MPLS): fundamentos y aplicaciones», 2006.
- [26] E. C. Rosen y Y. Rekhter, «Bgp/mppls vpns», 1999.
- [27] F. Palmieri, «VPN scalability over high performance backbones evaluating MPLS VPN against traditional approaches», en *Computers and Communication, 2003.(ISCC 2003). Proceedings. Eighth IEEE International Symposium on*, 2003, pp. 975–981.
- [28] J.-T. Park, M.-H. Kwon, y J.-H. Kwon, «Fast restoration of resilience-guaranteed segments under multiple link failures in a general mesh-type MPLS/GMPLS network», en *Information Networking. Advances in Data Communications and Wireless Networks*, Springer, 2006, pp. 359–368.
- [29] E. Mannie, «Generalized multi-protocol label switching (GMPLS) architecture», *Interface*, vol. 501, p. 19, 2004.
- [30] L. Berger, «Generalized multi-protocol label switching (GMPLS) signaling functional description», 2003.
- [31] E. Oki, K. Shiimoto, D. Shimazaki, N. Yamanaka, W. Imajuku, y Y. Takigawa, «Dynamic multilayer routing schemes in GMPLS-based IP+ optical networks», *Communications Magazine, IEEE*, vol. 43, n.º 1, pp. 108–114, 2005.
- [32] K. Sato, N. Yamanaka, Y. Takigawa, M. Koga, S. Okamoto, K. Shiimoto, E. Oki, y W. Imajuku, «GMPLS-based photonic multilayer router (Hikari router) architecture: an overview of traffic engineering and signaling technology», *Communications Magazine, IEEE*, vol. 40, n.º 3, pp. 96–101, 2002.
- [33] L. Berger, «Generalized multi-protocol label switching (GMPLS) signaling resource reservation protocol-traffic engineering (RSVP-TE) extensions», 2003.

- [34] K. Kompella y Y. Rekhter, «Label switched paths (LSP) hierarchy with generalized multi-protocol label switching (GMPLS) traffic engineering (TE)», 2005.
- [35] D. O. Awduche, «Applicability statement for extensions to RSVP for LSP-tunnels», 2001.
- [36] Ó. Calderón y X. Hesselbach, «Consideraciones sobre el balanceo de carga en redes MPLS», XIX Simposium Nacional de la Unión Científica Internacional de Radio, Barcelona, 2004.
- [37] S. Yoon, H. Lee, D. Choi, Y. Kim, G. Lee, y M. Lee, «An efficient recovery mechanism for MPLS-based protection LSP», 2001, pp. 75–79.
- [38] J. Ash, Y. Lee, P. Ashwood-Smith, B. Jamoussi, D. Fedyk, D. Skalecki, y L. Li, «LSP Modification Using CR-LDP», RFC3214 (Jan. 2002), 2002.
- [39] J.-L. Chen, S.-W. Liu, S.-L. Wu, y M.-C. Chen, «Cross-layer and cognitive QoS management system for next-generation networking», *International Journal of Communication Systems*, vol. 24, n.º 9, pp. 1150–1162, 2011.
- [40] D. Blandón, Y. Díaz, J. C. Cuellar, C. Ochoa, F. G. Guerrero, y A. Navarro, «Medición de la calidad del servicio.», *INTERACTIC*, vol. 01, p. 16, oct. 2008.
- [41] R. I. Davis, S. Kollmann, V. Pollex, y F. Slomka, «Controller area network (can) schedulability analysis with fifo queues», en *Real-Time Systems (ECRTS)*, 2011 23rd Euromicro Conference on, 2011, pp. 45–56.
- [42] R. I. Davis, S. Kollmann, V. Pollex, y F. Slomka, «Schedulability analysis for Controller Area Network (CAN) with FIFO queues priority queues and gateways», *Real-Time Systems*, vol. 49, n.º 1, pp. 73–116, 2013.
- [43] L. Shankar y S. Ambe, Weighted fair queuing (WFQ) shaper. Google Patents, 2003.
- [44] J. C. Bennett y H. Zhang, «Why WFQ is not good enough for integrated services networks», en *Proceedings of NOSSDAV*, 1996, vol. 96, pp. 524–532.
- [45] V. A. Siris y G. Stamatakis, «A dynamic CBWFQ scheme for service differentiation in WLANs», en *Personal, Indoor and Mobile Radio Communications*, 2005. PIMRC 2005. IEEE 16th International Symposium on, 2005, vol. 4, pp. 2665–2669.
- [46] A. Asghari, S. Cheng, J. D. Kirby, D. Mangiardi, y J. E. Schaff, Priority queuing in a next generation communication network. Google Patents, 2014.
- [47] J. C. Cuéllar Quiñones, «Análisis de configuraciones en el núcleo de una red NGN para garantizar QoS», *Sistemas y Telemática*, 2010.
- [48] O. Modeler, «OPNET Modeler manuals», *MIL*, vol. 3, pp. 1989–1997, 2007.

GLOSARIO

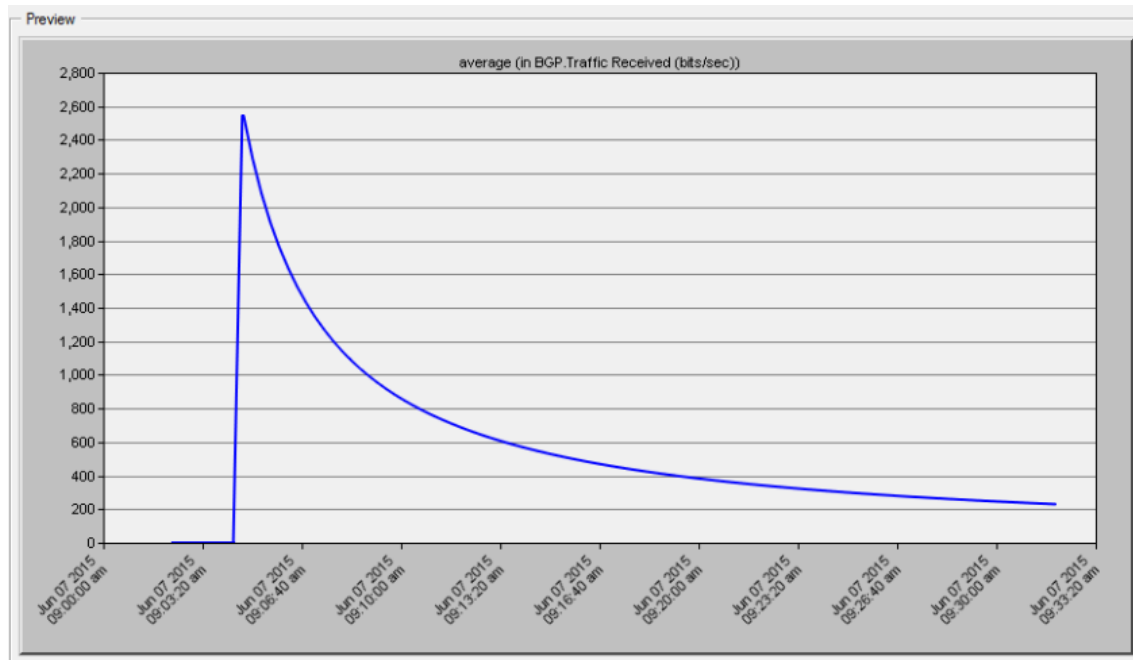
NGN (Redes de Próxima Generación).
QoS (Calidad de Servicio).
TE (Ingeniería de Trafico).
MPLS (Multiprotocol Label Switching).
ITU (Unión Internacional de Telecomunicaciones).
IP (Internet Protocol).
ATM (Asynchronous Transfer Mode).
MGs (Media Gateways).
AMGs (Access Media Gateways).
TMGs (Trunk Media Gateways).
SMGs (Signaling Media Gateways).
IADs (Integrate Access Devices).
UMG (Universal Media Gateways).
CRS (Cell Switching Router).
IETF (Grupo Especial sobre Ingeniería de Internet).
ARIS (Aggregate Route-based IP Switching).
ASIC (Circuito Integrado para Aplicaciones Específicas).
RSVP (Protocolo de Reserva de Recursos).
LAN (Local Area Network).
NSP (Network Service Provider).
LSP (Label Switched Path).
LSR (Label Switching Router).
OSPF (Open Shortest Path First).
Diff-Serv (Differentiated Services).
VPN (Virtual Private Network).
PPP (Point-to-Point Protocol).
SDH (Synchronous Digital Hierarchy).
IPX (Internetwork Packet Exchange).

OSI (Open System Interconnection).
WAN (Wide Area Network).
IGP (Interior Gateway Protocol).
IS-IS (Intermediate System to Intermediate System).
RIP (Routing Information Protocol).
LDP (Label Distribution Protocol).
CBR (Constraint-based Routing).
CoS (Clases de Servicio).
www (World Wide Web).
ToS (Type of Service).
EXP (Exptime).
PVC (Circuitos Privados Virtuales).
GMPLS (Multiprotocol Label Switching Generalized).
SONET (Synchronous Optical Network).
SDH (Synchronous Digital Hierarchy).
PDH (Plesiochronous Digital Hierarchy).
DWDM (Multiplexación por División en Longitudes de Onda).
TDM (Time Division Multiplex).
PSC (Packet Switch Capable).
L2SC (Layer-2 Switch Capable).
LSC (Lambda Switch Capable).
FSC (Fiber-Switch Capable).
SPF (Sender Police Framework).
AS (Sistema Autónomos).
RD (Routing Domain).
LMP (Link Management Protocol).
LP (Labeled Packets).
LER (Enrutadores de Etiqueta de Borde).
FEC (Forward Equivalence Class).
PHP (Hypertext Pre-Processor).
LDP (Label Distribution Protocol).

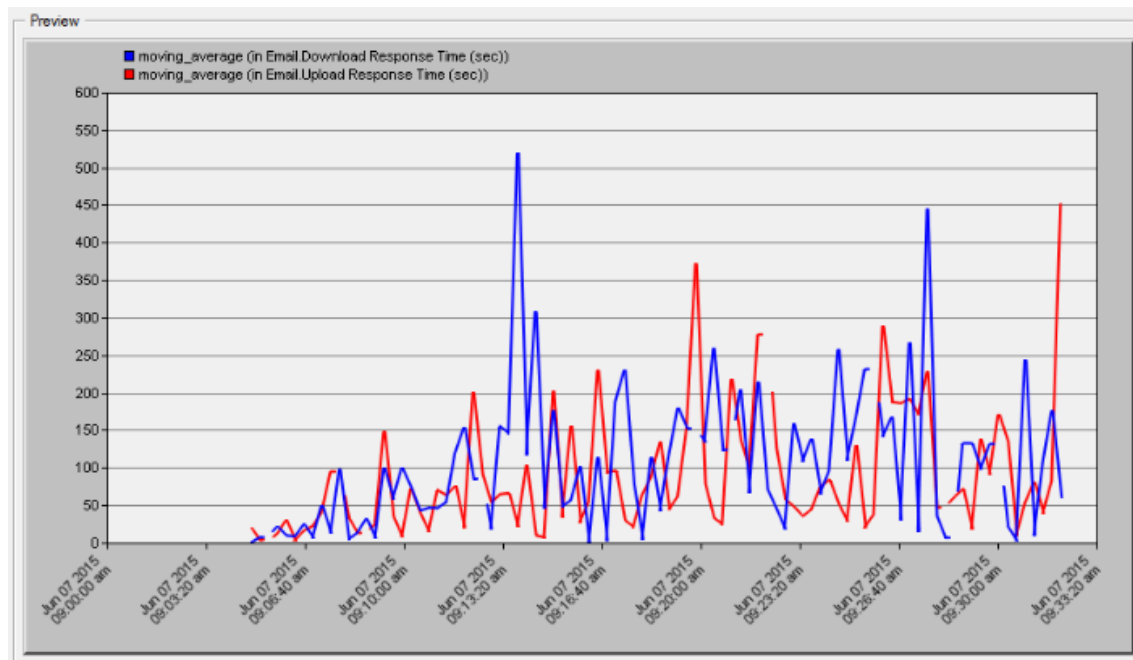
TCP (Transmission Control Protocol).
UDP (User Datagram Protocol).
ICMP (Internet Control Message Protocol).
DSCP (Differentiated Services Code Point).
LIB (Label Information Base).
FIFO (First-In First-Out).
WFQ (Weight Fair Queuing).
GPS (Generalized Processor Sharing).
FQ (Fair Queueing).
ACL (Access Control List).
PQ (Priority Queueing).
CQ (Custom Queueing).
PE (Provider Edge).
CE (Customer Edge).
DES (Simulación de Eventos Discretos).
OOP (Programación Orientada a Objetos).
FTP (File Transfer Protocol).
HTTP (Hypertext Transfer Protocol).
VRF (Virtual Routing and Forwarding).
LSE (Link State Algorithm).
LSDB (Link-State Database).
ToS (Tipo de Servicio).

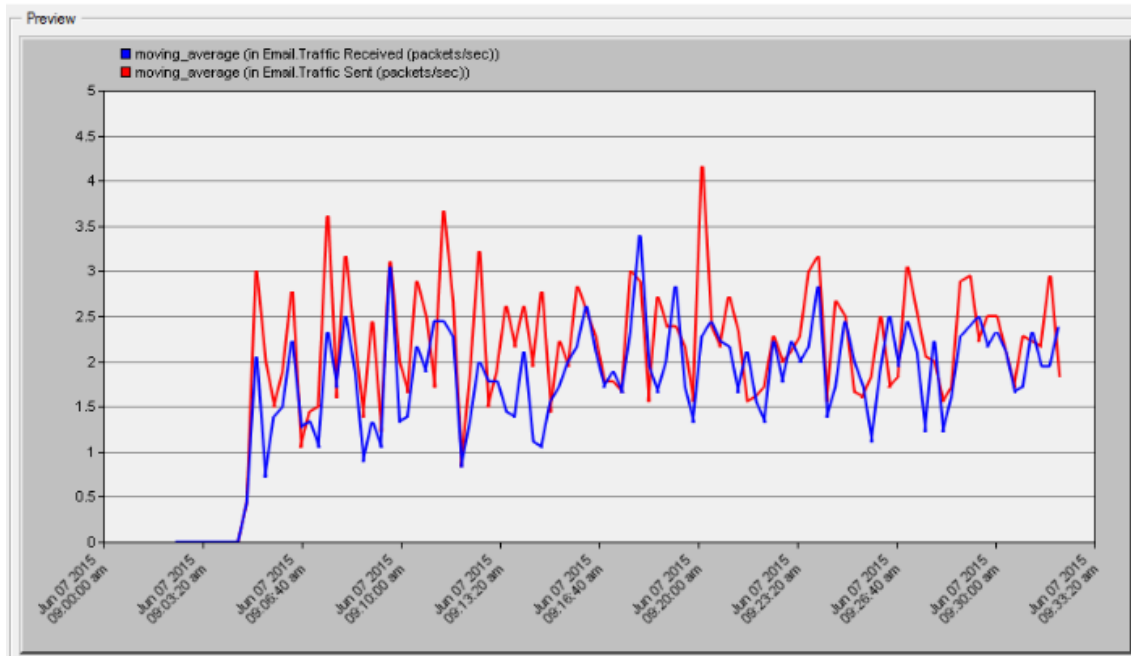
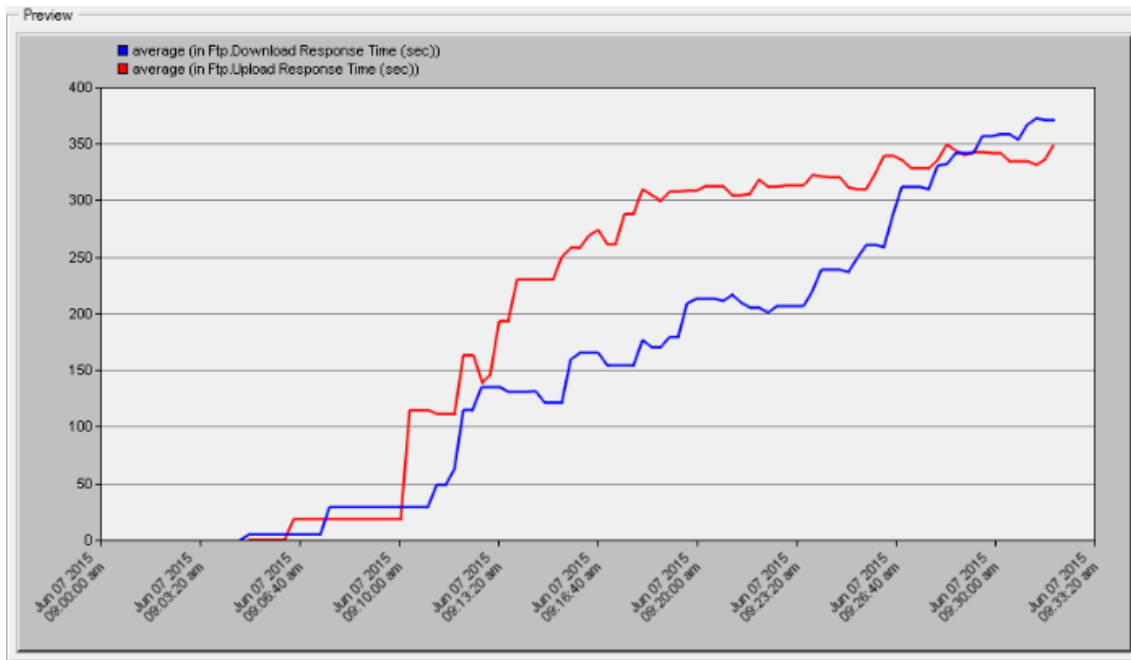
ANEXOS

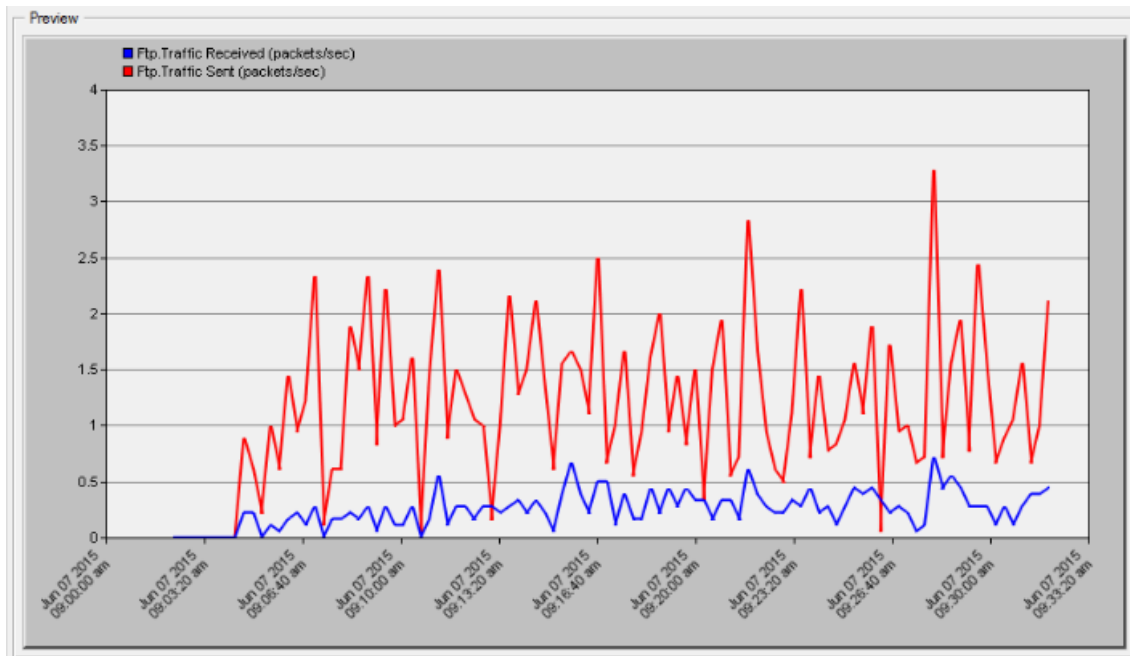
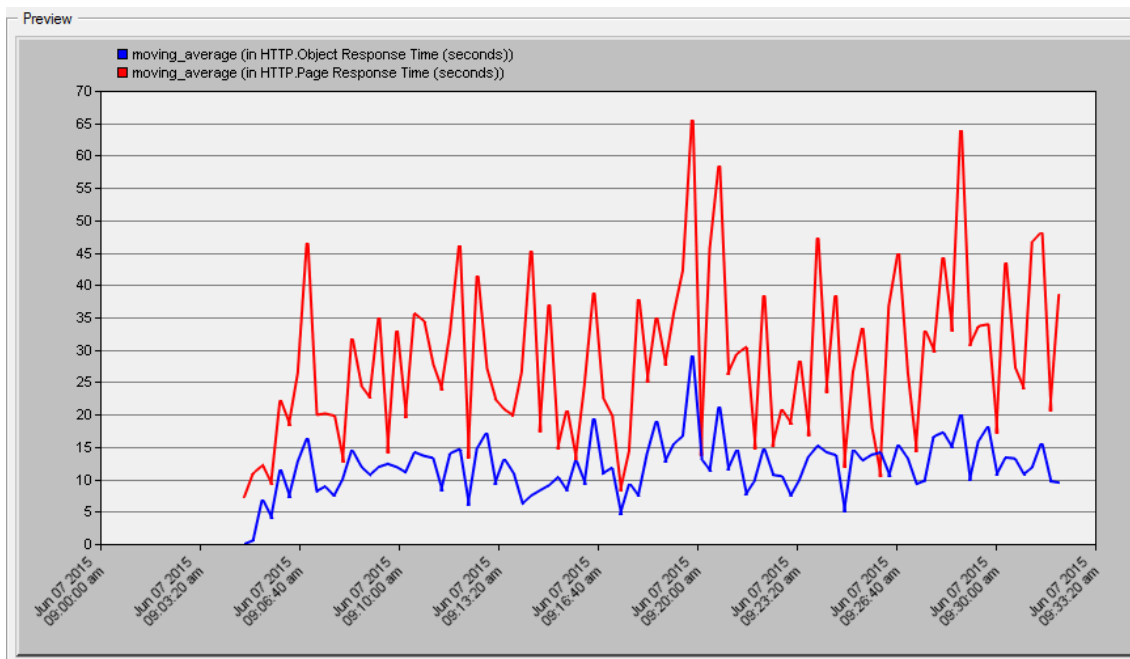
Anexo 1: Tráfico recibido para BGP.

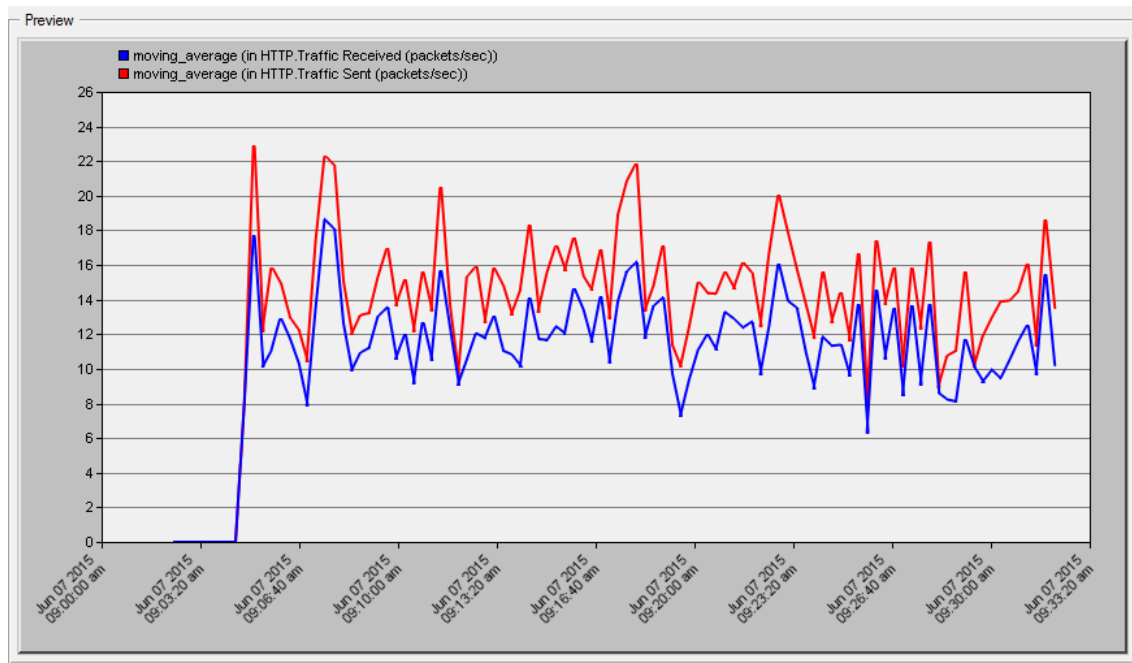
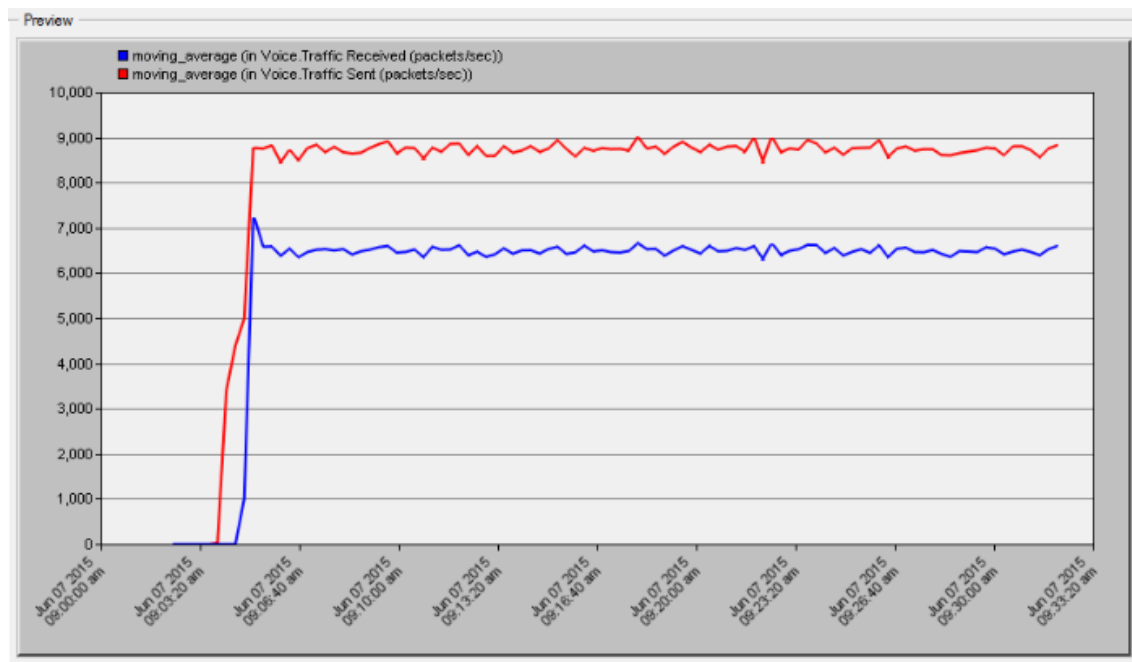


Anexo 2: Tiempo de Respuesta para descarga y subida de archivos Email.

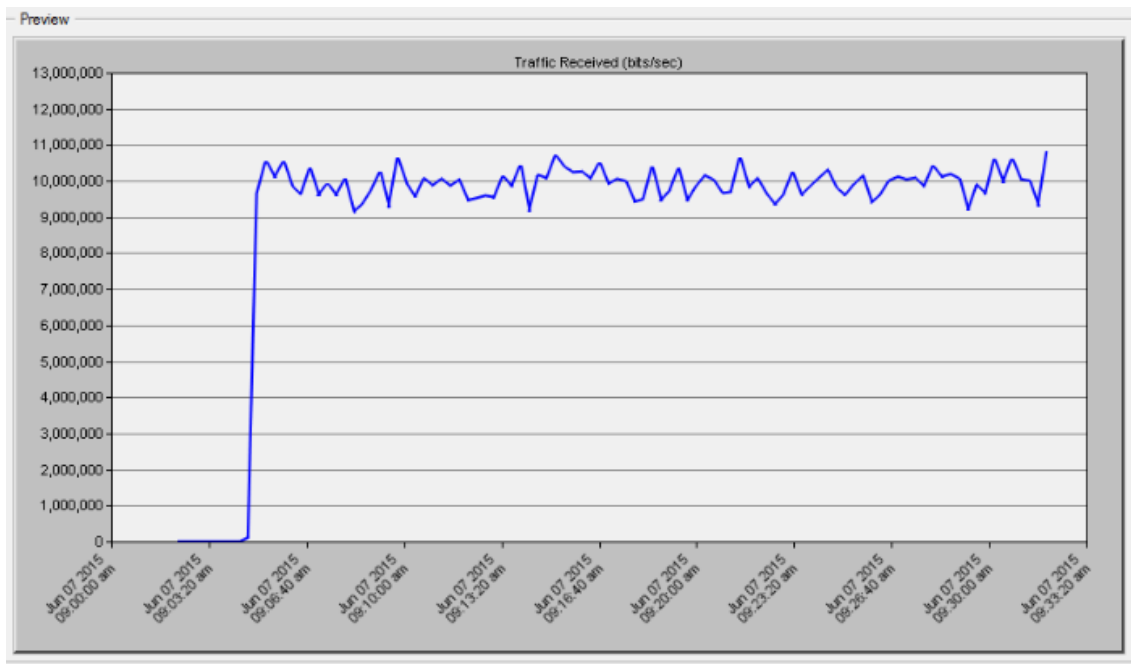


Anexo 3: Tráfico enviado y recibido para aplicaciones Email.**Anexo 4: Tiempo de Respuesta para descarga y subida de archivos FTP.**

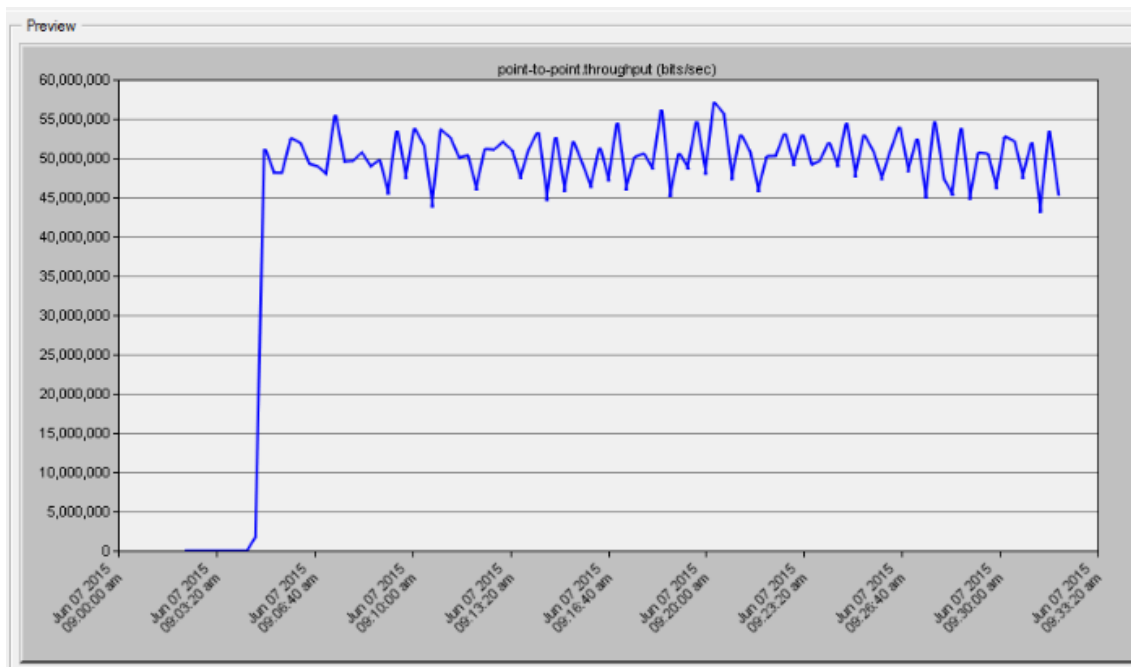
Anexo 5: Tráfico recibido y enviado para la aplicación FTP.**Anexo 6: Tiempo de respuesta para un objeto y una página utilizando el protocolo HTTP.**

Anexo 7: Tráfico recibido y enviado para la aplicación HTTP.**Anexo 8: Tráfico recibido y enviado para aplicaciones de Voz.**

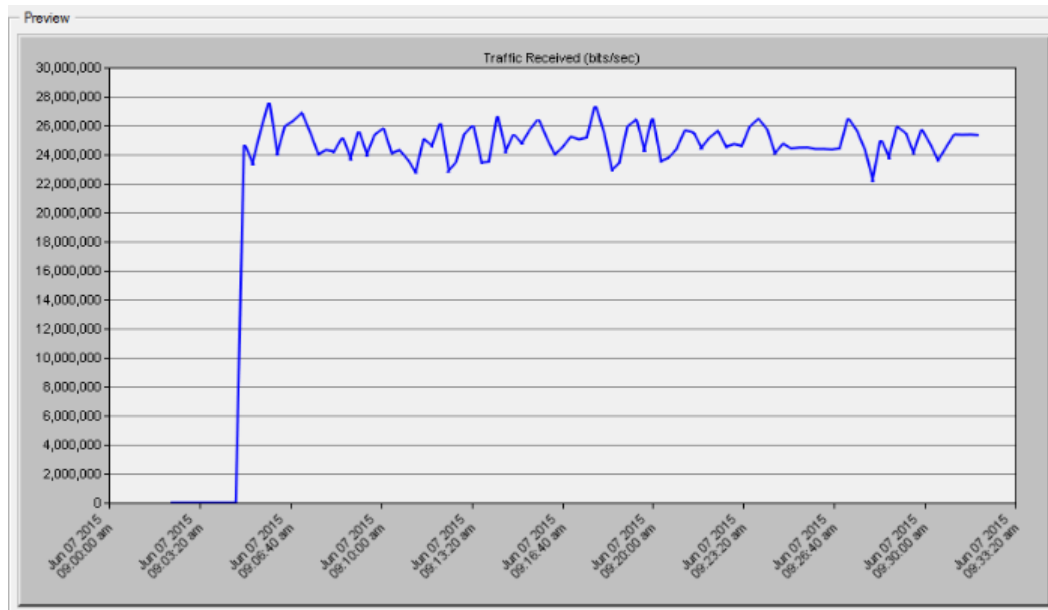
Anexo 9: Razón de Transferencia para las demandas simulando el tráfico de videoconferencia.



Anexo 10: Razón de Transferencia punto a punto para tráfico de voz existente entre los softswitch y la red NGN Villa Clara.



Anexo 11: Razón de Transferencia punto a punto para tráfico de voz existente entre los softswitch Las Tunas y Vedado.



Anexo 12: Perfiles creados en la Red.

(Profiles) Attributes

Type: Utilities

Attribute	Value
Profile Name	Profiles
Profile Configuration	(...)

(Profile Configuration) Table

Profile Name	Applications	Operation Mode	Start Time (seconds)	Duration (seconds)
Voz	(...)	Simultaneous	uniform (100,110)	End of Simulation
Video	(...)	Simultaneous	uniform (100,110)	End of Simulation
Datos	(...)	Simultaneous	uniform (100,110)	End of Simulation

(Applications) Table

Name	Start Time Offset (seconds)	Duration (seconds)	Repeatability
... Call (PCM Quality)	Voice over IP Call ...	uniform (5,10)	End of Profile

1 Rows Delete Insert Duplicate Move Up Move Down

Details Promote ☒ Show row labels OK Cancel

Anexo 13: Tabla VRF de Audio para el ISAM de Santa Clara.

	Destination	Source Protocol	Route Preference	Metric	Next Hop Address	Next Hop Node	Outgoing Interface	Outgoing LSP	Insertion Time (secs)	
1	192.0.46.0/24	IBGP	200	0	192.0.73.1	subnet_0_Corrillo	IF0	SANTA CLARA->Corrillo	160.007	
2	192.0.48.0/24	IBGP	200	0	192.0.83.1	subnet_0_Quemado	IF0	SANTA CLARA->Quemado	160.011	
3	192.0.50.0/24	IBGP	200	0	192.0.72.1	subnet_0_Sagua	IF0	SANTA CLARA->Sagua	160.008	
4	192.0.52.0/24	IBGP	200	0	192.0.75.1	subnet_0_Ciuentes	IF0	SANTA CLARA->Ciuentes	160.010	
5	192.0.54.0/24	IBGP	200	0	192.0.84.1	subnet_0_Encrucijada	IF0	SANTA CLARA->Encrucijada	160.004	
6	192.0.56.0/24	IBGP	200	0	192.0.74.1	subnet_0_Santo Domingo	IF0	SANTA CLARA->Santo Domingo	160.007	
7	192.0.58.0/24	IBGP	200	0	192.0.76.1	subnet_0_Ranchuelo	IF0	SANTA CLARA->Ranchuelo	160.011	
8	192.0.60.0/24	IBGP	200	0	192.0.78.1	subnet_0_Manicaragua	IF0	SANTA CLARA->Manicaragua	160.003	
9	192.0.63.0/24	Direct	0	0	192.0.63.2	subnet_0_SANTA CLARA	IF8	N/A	0.000	
10	192.0.65.0/24	IBGP	200	0	192.0.79.1	subnet_0_Placetas	IF0	SANTA CLARA->Placetas	160.003	
11	192.0.66.0/24	IBGP	200	0	192.0.80.1	subnet_0_Camajuari	IF0	SANTA CLARA->Camajuari	160.003	
12	192.0.68.0/24	IBGP	200	0	192.0.81.1	subnet_0_Remedios	IF0	SANTA CLARA->Remedios	160.003	
13	192.0.71.0/24	IBGP	200	0	192.0.82.1	subnet_0_Cabanién	IF0	SANTA CLARA->Cabanién	160.003	
14										
15	Gateway of last resort is not set									
16										

Anexo 14: Red simulada y modelada para Villa Clara.

