

**ESTADO DEL ARTE EN LA APLICACIÓN DE INGENIERIA DE TRAFICO EN
REDES IP**

**ING. LUIS ALFREDO ALVAREZ ESCALANTE
ING. EDISON QUINTERO MARIN**

**UNIVERSIDAD PONTIFICIA BOLIVARIANA - UPB
FACULTAD DE INGENIERIA DE ELECTRONICA
ESPECIALIZACION EN TELECOMUNICACIONES
BUCARAMANGA
2010**

**ESTADO DEL ARTE EN LA APLICACIÓN DE INGENIERIA DE TRAFICO EN
REDES IP**

**ING. LUIS ALFREDO ALVAREZ ESCALANTE
ING. EDISON QUINTERO MARIN**

Monografía para optar el título de Especialista en Telecomunicaciones

**Director
PhD. Jhon Jairo Padilla Aguilar**

**UNIVERSIDAD PONTIFICIA BOLIVARIANA – UPB
FACULTAD DE INGENIERIA ELECTRONICA
ESPECIALIZACION EN TELECOMUNICACIONES
BUCARAMANGA
2010**

Nota de Aceptación

Presidente del Jurado

Jurado

Jurado

Bucaramanga, Abril 10 de 2010

A nuestros Padres, hermanos y amigos quienes con su ayuda incondicional nos apoyaron permanentemente para poder alcanzar el título de especialistas en Telecomunicaciones.

(Dedicatoria)

AGRADECIMIENTOS

A nuestro distinguido director:

Quien siempre estuvo orientando nuestro trabajo con su desinteresada y oportuna labor de transmisión del saber, así como sus acertados consejos y sugerencias.

A nuestros amigos y compañeros:

Quienes con su ayuda y fraternal compañerismo siempre estuvieron dispuestos a colaborarnos y apoyarnos.

Un agradecimiento especial para cada uno de los docentes de la Especialización en Telecomunicaciones de la Universidad Pontificia Bolivariana seccional Bucaramanga, por su entrega y profesionalismo en sus labores académicas. Así como a todos a aquellos que de una u otra forma intervinieron en nuestra formación. ¡Gracias!

TABLA DE CONTENIDO

	Pág.
DEDICATORIA	
AGRADECIMIENTOS	
LISTA DE GRAFICAS	
INTRODUCCION	15
MARCO TEORICO	
1. MECANISMOS ESTANDARIZADOS PARA GARANTIZAR CALIDAD DE SERVICIO EN LAS REDES IP EN LA ACTUALIDAD	21
1.1 MODELO INTSERV / RSVP (RESOURCE RESERVATION PROTOCOL)	21
1.1.1. Tipos de servicio en IntServ.	23
1.1.2. Características del Protocolo RSVP.	24
1.1.3. Componentes del protocolo RSVP.	24
1.1.4. Desventajas del protocolo RSVP.	25
1.2 MODELO DIFFSERV (DIFFERENTIATED SERVICES – DIFERENCIACIÓN DE SERVICIOS)	25
1.3 COMBINACIÓN DE RSVP Y DIFFSERV	29
2. MANEJO DE ETIQUETAS EN LAS REDES IP – (Protocol LDP)	31
3. PROTOCOLOS DE ENRUTAMIENTO ACTUALES EN REDES IP	32

3.1. PROTOCOLO OSPF (OPEN SHORT PATH FIRST - ABRIR PRIMERO LA RUTA MÁS CORTA)	32
3.1.1. Tipos de mensajes del Protocolo OSPF.	32
3.1.2. Funcionamiento básico de OSPF.	33
3.1.3. Mapa de Red Local.	33
3.1.4. Características de OSPF.	34
3.1.5. Relación con los vecinos en OSPF.	37
3.2. BORDER GATEWAY PROTOCOL (BGP) – PROTOCOLO DE GATEWAY FRONTERIZO	39
3.3. ROUTING INFORMATION PROTOCOL RIP – PROTOCOLO DE ENCAMINAMIENTO DE INFORMACIÓN.	40
3.4. INTERMEDIATE SYSTEM TO INTERMEDIATE SYSTEM (IS-IS)	41
4. MULTIPROTOCOL LABEL SWITCHING MPLS – CONMUTACIÓN DE ETIQUETAS MULTIPROTOCOLO	44
4.1. CARACTERÍSTICAS DE LAS REDES MPLS	46
4.2. ROUTING EN MPLS	49
4.3. FUNCIONAMIENTO DE UNA RED MPLS	50
5. PROTOCOLOS COMPLEMENTARIOS PARA HACER INGENIERÍA DE TRÁFICO EN REDES MPLS	51

5.1. EXTENSIÓN DE INGENIERÍA DE TRÁFICO PARA PROTOCOLOS DE ENRUTAMIENTO	52
5.2. PAQUETES ETIQUETADOS Y LSP	53
5.3. DISTRIBUCIÓN DE ETIQUETAS	56
5.4. CR-LDP (CONSTRAINT-BASED ROUTING – LABEL DISTRIBUTION PROTOCOL)	56
5.5. RSVP-TE (RESOURCE RESERVATION PROTOCOL– TRAFFIC ENGINEERING)	59
DESCRIPCION DEL PROBLEMA	
6. PORQUÉ Y COMO APLICAR INGENIERIA DE TRÁFICO A LAS REDES IP	63
6.1. QUÉ ES EL TRÁFICO EN REDES IP?	64
6.2. TRÁFICO Y MEDIDAS DE DESEMPEÑO	64
6.3. CARACTERIZACIÓN DEL TRÁFICO	65
6.4. ACERCA DE LAS CARACTERÍSTICAS DEL COMPORTAMIENTO ALEATORIO	66
6.5. CUÁLES SON LOS PROBLEMAS DE UNA RED CUANDO EXISTE MÁS DE UN ENLACE?	66
6.6. OPTIMIZACIÓN DEL FLUJO DE RED	68

6.7. ENRUTAMIENTO POR EL PRIMER CAMINO MÁS	72
--	----

CORTO Y FLUJO DE RED

PROPUESTAS EXISTENTES

7. INGENIERIA DE TRAFICO CON PROTOCOLOS DE ENRUTAMIENTO IP TRADICIONALES	80
--	----

7.1. VENTAJAS DE LA UTILIZACIÓN TRADICIONAL DEL OSPF / IS-IS	82
--	----

7.2. SELECCIÓN DE CAMINO BASADO EN CONFIGURACIÓN IGP	83
--	----

7.3. CONTROL: RECONFIGURACIÓN DE PESOS IGP	84
--	----

7.4. PROPIEDADES DE RENDIMIENTO	84
---------------------------------	----

7.5. TOPOLOGÍAS FIJAS Y DEMANDAS DE TRÁFICO	85
---	----

7.6. COMPARACIÓN DEL RENDIMIENTO CON MÁXIMA UTILIZACIÓN DEL ENLACE	85
--	----

7.7. COMPARACIÓN DE RENDIMIENTO CON UNA RED DE GRAN CAPACIDAD	86
---	----

7.8. REDUCIENDO EL NÚMERO DE LOS VALORES DE PESO	88
--	----

7.9. CAMBIANDO LA DEMANDA DE TRÁFICO	88
--------------------------------------	----

7.10.	POCOS CAMBIOS PARA LOS PESOS DE ENLACE	89
8.	HACIA LA BAJA COMPLEJIDAD DE LA INGENIERÍA DE TRÁFICO EN INTERNET: LA ADAPTACIÓN DEL ALGORITMO MULTI CAMINO	90
8.1.	INGENIERÍA DE TRÁFICO PARA ENRUTAMIENTO CON ESTADO DE ENLACE	90
8.2.	AVANCES RECIENTES EN ENRUTAMIENTO DE MÚLTIPLES CAMINOS	91
8.3.	ALGORITMO DE ENRUTAMIENTO ADAPTATIVO MULTICAMINO (AMP)	92
8.4.	CÁLCULO DE MÚLTIPLES CAMINOS	94
8.5.	MÉTRICAS DE CARGA DE ENLACE – CÁLCULO DE LA CARGA EN EL ENLACE	95
8.6.	SEÑALIZACIÓN PARA ENRUTAMIENTO ADAPTATIVO MULTICAMINO (AMP)	96
8.7.	REENVÍO DE PAQUETES	97
8.8.	BALANCEO DE CARGA (AMP)	98

CONCLUSIONES

BIBLIOGRAFIA

INDICE DE GRAFICAS

	Pág
Gráfica 1: Enrutamiento entre sistemas autónomos	18
Gráfica 2: Modelo IntServ / RSVP	22
Gráfica 3: Modelo DiffServ	25
Gráfica 4: Cabecera IPv4 antes de DiffServ	27
Gráfica 5: Cabecera IPv4 con DiffServ	27
Gráfica 6: Campo TOS	28
Gráfica 7: Campo DS	28
Gráfica 8: Implementación de DiffServ en Routers	29
Gráfica 9: Combinación de RSVP y DiffServ	29
Gráfica 10: Áreas OSPF	36
Gráfica 11: Etiqueta MPLS	47
Gráfica 12: Formato de etiqueta MPLS	47
Gráfica 13: Ubicación de etiqueta según tecnología	48
Gráfica 14: Enrutamiento MPLS	49
Gráfica 15: Intercambio de etiquetas en MPLS	49
Gráfica 16: Caminos de etiquetas conmutadas usando un túnel LSP	54
Gráfica 17: Intercambio y conmutación de etiquetas entre caminos	55
Gráfica 18: Proceso para determinar ruta y reserva de recursos	58
Gráfica 19: Establecimiento de un LSP válido	59

Gráfica 20: Distribución de etiquetas usando RSVP-TE	60
Gráfica 21: Ingeniería de tráfico IP – Marco Arquitectónico	67
Gráfica 22: Ejemplo para red de 4 nodos	68
Gráfica 23: Red con 4 nodos y la misma capacidad de los enlaces	70
Gráfica 24: Red con 4 nodos y diferente capacidad en los enlaces	71
Gráfica 25: Red de 4 nodos con 5 enlaces	76
Gráfica 26: Red de 4 nodos con pesos de enlace	78
Gráfica 27: Enrutamiento del camino más corto dentro de un Sistema Autónomo basado sobre OSPF/IS-IS con pesos de enlace	80
Gráfica 28: Enrutar las mismas demandas con diferentes configuraciones de pesos de enlace	82
Gráfica 29: Costo del enlace en función de la carga y con capacidad 1	86
Gráfica 30: Costo de Red con Mayor Capacidad Vrs Demanda para una propuesta de Backbone AT&T	87
Gráfica 31: Idea Básica del Mecanismo Contrapresión	93
Gráfica 32: Ejemplo para un mensaje contrapresión (BMs- backpressure messages)	97
Gráfica 33: Promedio de tiempo de respuesta de la página Web de AT&T en EE.UU	99

RESUMEN GENERAL DE TRABAJO DE GRADO

TITULO: ESTADO DEL ARTE EN LA APLICACIÓN DE INGENIERIA DE TRAFICO EN LAS REDES IP
AUTORES: LUIS ALFREDO ALVAREZ ESCALANTE
EDISON QUINTERO MARIN
FACULTAD: ESPECIALIZACION EN TELECOMUNICACIONES
DIRECTOR: JHON JAIRO PADILLA AGUILAR

RESUMEN

Actualmente las redes IP presentan un gran crecimiento en la información a transmitir debido al mayor número de aplicaciones existentes (Videostream, VozIP, Audiostream), las cuales integran grandes flujos de información (audio, voz, datos y video simultáneamente). De tal manera que la demanda por estos nuevos servicios requiere que estas redes sean lo más eficientes posibles en la administración de sus recursos y así poder garantizar calidad en el servicio. Es por ello que nace el concepto de ingeniería de tráfico que aborda los objetivos de una red IP, dimensionándolos en el tiempo según el crecimiento de la misma en usuarios, volumen de tráfico así como en el comportamiento general de la red. Además no se refiere a la adición de nueva capacidad en los enlaces, sino que hace referencia a optimizar los recursos ya existentes. La ingeniería de tráfico tiene como finalidad obtener el sistema óptimo de pesos en los enlaces al tiempo que reconoce que la red utiliza el enrutamiento del camino más corto para la transmisión. Para ello se debe conocer de antemano la matriz de tráfico, la topología de la red y su configuración. A través de esta monografía se analiza cómo se implementa la ingeniería de tráfico en las redes IP teniendo en cuenta que estas redes trabajan con protocolos de enrutamiento llamados protocolos de estado de enlace conocidos con los nombres de OSPF (Open Short Path First) / IS-IS (Intermediate System to Intermediate System), los cuales tiene como función el balancear la carga y escoger el camino más corto entre los nodos origen y destino. Para ello mediante los pesos de enlace se debe equilibrar el volumen de tráfico a transmitir en la red, teniendo en cuenta cómo funcionan dichas redes. Una vez se comprende como es el enrutamiento de estas redes en la actualidad, el siguiente paso es poder optimizarlas al máximo aplicando ingeniería de tráfico y para ello se analizaron varias propuestas al respecto. Una de ellas confirma que con estos protocolos de enrutamiento se obtiene una buena configuración de pesos en los enlaces para así satisfacer un rendimiento en particular de la red. Otras investigaciones afirman que estos enfoques no adaptan la ruta o camino a las condiciones actuales de tráfico en las redes. Ante esto surge la propuesta del algoritmo de Enrutamiento Adaptativo Multicamino – AMP, el cual es compatible con la arquitectura de enrutamiento existente. Dicho algoritmo tiene como objetivo permitir el balanceo de carga continua aplicando ingeniería de tráfico de forma autónoma en cada nodo. Ofrece la opción de utilizar el mejor camino encontrado logrando evitar la congestión en la transmisión de los datos ante cualquier cambio en el tráfico.

PALABRAS CLAVES: REDES IP, INGENIERIA DE TRAFICO, OSPF, IS-IS, PROTOCOLOS DE ENRUTAMIENTO IP

GENERAL SUMMARY OF WORK OF DEGREE

TITLE: STATE OF THE ART IN THE IMPLEMENTATION OF
TRAFFIC ENGINEERING IN IP NETWORKS
AUTHOR: LUIS ALFREDO ALVAREZ ESCALANTE
EDISON QUINTERO MARIN
FACULTY: SPECIALIZATION IN TELECOMMUNICATIONS
DIRECTOR: JHON JAIR PADILLA AGUILAR

ABSTRACT

Currently, IP networks have a tremendous growth in the information to be transmitted by the greater number of existing applications (Videostream, VoIP, audiostream), which comprise large flows of information (audio, voice, data and video simultaneously). So the demand for these new services requires that these networks are as efficient as possible in managing their resources and thereby ensure quality service. That is the reason because the concept of traffic engineering was created. It addresses aspects such as the objectives of an IP network, dimensions in time according to the same growth in users, traffic patterns and general behavior of the network. However, Traffic Engineering does not focus in to the addition of new capacity on the links, but refers to optimize existing resources. The traffic engineering aims to obtain the optimal weights on the links while, at the same time, it acknowledges that the network uses shortest path routing for transmission. With this purpose, the traffic matrix, network topology and configuration should be known in advance. Through this paper we analyze how to implement traffic engineering in IP networks taking into account that these networks work with routing protocols called link-state protocols, they are known by the names of OSPF (Open Short Path First) / IS- IS (Intermediate System to Intermediate System). Also, they have known which is the load balancing function and they should choose the shortest path between source and destination nodes. To do this by using the link weights, the network should balance the volume of traffic transmitted on the network. The next step is to optimize the data traffic by applying several proposals that are discussed in this document. One of them confirms that these routing protocols could obtained a good configuration of weights on the links in order to satisfy a particular performance of the network. Other studies argue that these approaches do not fit the route or path for current traffic conditions in the nets. Against this proposal, arises the Adaptive multipath routing algorithm - AMP, which is compatible with the existing routing architecture. This algorithm is intended to allow load balancing and traffic engineering to apply independently in each node. It offers the option of using the best path found in such way that it avoids congestion in the transmission of data with any change in traffic.

KEYWORDS: IP NETWORKS, TRAFFIC ENGINEERING, OSPF, IS-IS, IP ROUTING PROTOCOLS

INTRODUCCION

En la actualidad los sistemas de telecomunicaciones hacen parte del diario vivir de las personas, están presentes en todas sus actividades cotidianas tanto a nivel personal como empresarial. Estos sistemas vienen evolucionando constantemente hacia redes denominadas de próxima generación NGN, también conocidas como redes IP trayendo consigo una serie de nuevos conceptos que deben ser analizados y estudiados con el fin de contribuir a la convergencia de datos, voz y video bajo una misma arquitectura de red.

Uno de los grandes problemas que se presenta al enviar información usando redes IP, es mantener la integridad de los datos y garantizar una comunicación fluida, tanto en velocidad como en efectividad, es por esto que si se logra brindar calidad del servicio QoS en las aplicaciones que son ofrecidas a los usuarios finales, se logrará mejorar el Grado de Satisfacción GoS de ellos hacia las diferentes aplicaciones soportadas por la red¹.

Con este proyecto se busca realizar un análisis del estado del arte de los diferentes procesos tecnológicos que permiten un buen flujo de la información a través de la red y la búsqueda del mejor camino (enrutamiento) y deducir cual podría ser el que más garantiza la aplicación del concepto de Calidad del Servicio QoS para redes IP.

Para ello se hace necesario aplicar Ingeniería de Tráfico a estas redes, “la cual tiene como objetivo diseñar sistemas con un costo mínimo y con una capacidad tal que cumpla con un grado de servicio predefinido satisfaciendo la demanda de tráfico a futuro”².

Las recomendaciones de la ITU-T International Telecommunication Union con respecto a los componentes de la ingeniería de tráfico son clasificadas en cuatro categorías:

¹FORTZ, Bernard, REXFORD Jennifer and THORUP Mikkell. Universite Catholique de Louvain. Traffic Engineering With Traditional IP Routing Protocols. De: ScienceDirect 2009.

² PADILLA, Jhon Jairo. Componentes de la Ingeniería de Tráfico.pdf. Universidad Pontificia Bolivariana. Bucaramanga. 2009.

Caracterización de la demanda de tráfico: con el fin de obtener modelos matemáticos que describan aproximadamente el comportamiento aleatorio del tráfico de los usuarios. Una vez se modela el tráfico generado por los usuarios es posible determinar el comportamiento del sistema, obteniendo los parámetros que son relevantes para ser modelados (retardo, ancho de banda, longitudes de colas, etc).

Objetivos de Grado de Servicio GoS: es un conjunto de características de los servicios que determinan el grado de satisfacción de un usuario en relación a un servicio, los cuales suelen ser probabilidad de bloqueo y retardo.

Controles de tráfico y dimensionamiento: hace referencia a las herramientas de la ingeniería de tráfico en base a los métodos de diseño y a los mecanismos de gestión y operación de la red. Los primeros buscan dimensionar la red para así tener suficientes recursos y cumplir con la demanda de tráfico de los usuarios (recursos físicos y lógicos), mientras los segundos son necesarios para cumplir con el control de tráfico requerido para asegurar que se cumplan los objetivos de GoS.

Monitoreo del rendimiento: una vez la red este dimensionada, hay situaciones de sobrecarga y fallos no considerados en el dimensionamiento. Por tanto se requiere tomar acciones inmediatas para contrarrestar tales situaciones. Como consecuencia de ello se reconfiguran elementos en la red, se reconfigura el enrutamiento y se hacen ajustes en los parámetros de control de tráfico.³

Existen redes MPLS⁴ (Multiprotocol Label Switching – Conmutación de Etiquetas Multiprotocolo) que actualmente implementan el concepto de Ingeniería de Tráfico de un manera tal, que facilita el enrutamiento de los datos y el flujo de información se realiza de manera eficiente alcanzando altos grados de desempeño hacia los usuarios. Para ello se vale de protocolos de señalización tales como el RSVP⁵ (Resource Reservation Protocol - Protocolo de Señalización de Reserva de

³ PADILLA, Ibid. Págs: 1-30.

⁴ ROSEN E, VISWANATHAN A. Multiprotocol Label Switching Architecture.2001.
[Http://tools.ietf.org/html/rfc3031](http://tools.ietf.org/html/rfc3031).

⁵ MANKIN A, BAKER F, BRADEN B. Resource ReSerVation Protocol. 1997.
[Http://tools.ietf.org/html/rfc2208](http://tools.ietf.org/html/rfc2208)

Recursos) y el LDP⁶ (Label Distribution Protocol – Protocolo de distribución de Etiquetas) ambos en sus versiones especializadas RSVP-TE⁷ (Resource Reservation Protocol – Traffic Engineering) y CR-LDP⁸ (Constraint – Route Label Distribution Protocol).

Inicialmente se analizará el protocolo **RSVP** (Resource Reservation Protocol - Protocolo de Señalización de Reserva de Recursos) y el protocolo **LDP** (Label Distribution Protocol - Protocolo para la distribución de etiquetas). Además de los Protocolos **IGP** (Interior Gateway Protocol - Protocolo de enrutamiento de Gateway Interior) y los Protocolos **EGP** (Exterior Gateway Protocol – Protocolo de Enrutamiento Exterior) que son los utilizados en la actualidad para enrutamiento en las redes IP. Los protocolos IGP son utilizados para el enrutamiento intradominio, es decir los utilizados dentro del sistema autónomo. Ejemplo: RIP⁹, OSPF¹⁰, etc. Mientras que los protocolos EGP son utilizados para el enrutamiento interdominio es decir entre sistemas autónomos. Ejemplo: BGP¹¹, EGP¹², etc.

“Se entiende por Sistema Autónomo a un conjunto de redes gestionadas por una administración común y que comparten una estrategia de encaminamiento común”¹³.

La grafica 1 muestra cómo interactúan estos dos últimos protocolos entre sistemas autónomos independientes.

⁶ ANDERSON L, DOOLAN P, FELDMAN L. LDP Specification. 2001.

[Http://tools.ietf.org/html/rfc3036](http://tools.ietf.org/html/rfc3036)

⁷ AWDUCHE D, BERGER L. RSVP-TE: Extensions to RSVP for LSP Tunnels. 2001. [Http://tools.ietf.org/html/rfc3209](http://tools.ietf.org/html/rfc3209)

⁸ ASH J, GIRISH M, GRAY E. Applicability Statement for CR-LDP. 2002. [Http://tools.ietf.org/html/rfc3213](http://tools.ietf.org/html/rfc3213)

⁹ MALKIN, G. RIP Version 2. 1998. [Http://tools.ietf.org/html/rfc2453](http://tools.ietf.org/html/rfc2453)

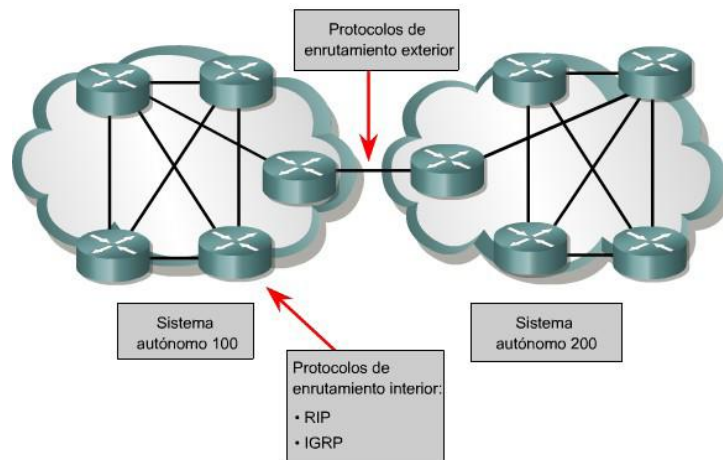
¹⁰ MOY, J. OSPF Version 2. 1998. [Http://tools.ietf.org/html/rfc2328](http://tools.ietf.org/html/rfc2328)

¹¹ REKHTER, Y. A Border Gateway Protocol. 1995. [Http://tools.ietf.org/html/rfc1771](http://tools.ietf.org/html/rfc1771)

¹² MILLS, D. Exterior Gateway Protocol Formal Specification. 1984. [Http://tools.ietf.org/html/rfc904](http://tools.ietf.org/html/rfc904)

¹³ HAWKINSON, J. Guidelines for Creation, Selection, and Registration of an Autonomous System (AS). 1996. [Http://tools.ietf.org/html/rfc1930](http://tools.ietf.org/html/rfc1930)

Gráfica 1
Enrutamiento entre Sistemas Autónomos



Tomada de: [http:// www.raap.org.pe/docs/raap2_RipOspf.pdf](http://www.raap.org.pe/docs/raap2_RipOspf.pdf)

Los protocolos IGP pueden a su vez clasificarse en Protocolos de Vector-Distancia ó de Estados de Enlace. Para efectos de la investigación se trabajará los protocolos de Estados de Enlace tales como:

- **OSPF** (Open Shortest Path First) – Primero la ruta libre mas corta.
- **IS-IS¹⁴** (Intermediate System to Intermediate System) - Sistema Intermedio a Sistema Intermedio.

Los protocolos de enrutamiento de Estado de Enlace fueron diseñados para superar las limitaciones de los protocolos de enrutamiento Vector Distancia. Estos últimos solo intercambian actualizaciones con sus vecinos inmediatos, mientras que los primeros tienen pleno conocimiento de los routers distantes y la forma como se interconectan, intercambiando información a través de un área más amplia; utilizan métricas de costo para seleccionar rutas a través de la red, utilizan inundación LSA para informar sobre cambios en la red convergiendo más rápidamente. En los protocolos Estado de Enlace cada router tiene una topología de su propia red,

¹⁴ ORAN D. OSI IS-IS Intra-domain Routing Protocol. 1990.
[Http://tools.ietf.org/html/rfc1142](http://tools.ietf.org/html/rfc1142)

consumen menos ancho de banda, y se ejecuta el algoritmo Dijkstra-Primer camino más corto. El protocolo OSPF organiza un sistema autónomo en áreas, donde dichas áreas son grupos lógicos de routers cuya información se puede resumir para el resto de la red¹⁵.

Todo este análisis tiene como propósito entender el funcionamiento del enrutamiento de las redes IP tanto a nivel interno como externo.

De cierto modo las redes IP se auto gestionan por medio de estos protocolos, haciendo que la comunicación sea de la mejor calidad posible. Para ello se manejan diferentes métricas ó características de las rutas como son: ancho de banda, retardo, carga, confiabilidad, número de saltos, costo, etc. Donde además ante cualquier cambio en la configuración de la topología de red, los routers están en la capacidad de actualizar sus tablas de enrutamiento para proceder a seleccionar nuevos caminos posibles. Sin embargo, estos mecanismos no garantizan que la red sea lo más eficiente posible. Pues puede pasar que un enlace esté congestionado a pesar que existan enlaces subutilizados en otras partes de la red ó un servicio requerido puede viajar sobre una ruta con alto retardo de propagación cuando un camino de baja latencia está disponible. Es por eso que se requiere hacer más eficiente la administración y el uso de los recursos que se encuentran disponibles en la red, mejorando los tiempos de respuesta hacia los usuarios.¹⁶

Aplicando Ingeniería de Tráfico a las actuales restricciones en la red, es posible lograr un mejor grado de servicio GoS y una mejor calidad del servicio QoS.

Para aclarar todos estos conceptos en esta monografía se consultaron referencias bibliográficas, se analizaron artículos publicados por investigadores internacionales y se revisaron bases de datos aplicadas a la ingeniería, de las cuales se extrajo información que permite perfeccionar el estado del arte al respecto y así fomentar futuras investigaciones que retomen el tema y puedan continuar con un proceso investigativo que facilite la solución a las nuevas exigencias de las redes IPv6.

¹⁵ MOY, J. Op. Cit. Págs: 34-40.

¹⁶ FORTZ, Bernard. Op. Cit Págs: 1-3.

PARTE I: MARCO TEORICO

1. MECANISMOS ESTANDARIZADOS PARA GARANTIZAR CALIDAD DE SERVICIO EN LAS REDES IP EN LA ACTUALIDAD

En estos momentos existen dos mecanismos estandarizados para garantizar calidad de servicio en las redes IP; tales mecanismos están enfocados en la reserva y prioridad de los recursos.

- **InterServ / RSVP** (Integrated Services – Servicios Integrados / Resource Reservation Protocol – Protocolo de reserva de recursos).

“Este mecanismo reserva la capacidad solicitada de recursos por un flujo en todos los routers del camino con ayuda del protocolo RSVP. El usuario solicita de antemano los recursos que necesita; cada router del trayecto ha de tomar nota y efectuar la reserva solicitada”¹⁷

- **DiffServ** (Differentiated Services).

“Este mecanismo intenta evitar los problemas de escalabilidad que plantea IntServ/RSVP. Se basa en el marcado de paquetes únicamente. No hay reserva de recursos por flujo, no hay protocolo de señalización, no hay información de estado en los routers. Las garantías de calidad de servicio no son tan severas como en IntServ pero en muchos casos se consideran suficientes”¹⁸.

Estos dos mecanismos pueden complementarse en la misma red IP.

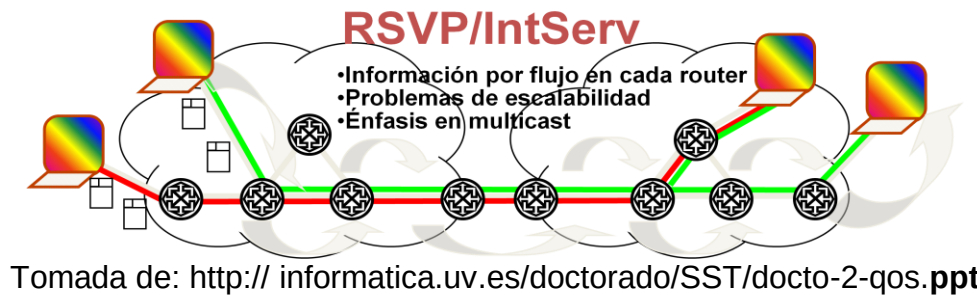
1.1. Modelo IntServ / RSVP (Resource Reservation Protocol)¹⁹

¹⁷ MANKIN A, BAKER F, BRADEN B. Op. Cit. Págs: 3-4.

¹⁸ BAKER, F. Definición de los Servicios de Campo diferenciada (DS Field) En el IPv4 y IPv6 Cabeceras. 1998. [Http://www.normes-Internet.com/normes.php?rfc=rfc2474&lang=es](http://www.normes-Internet.com/normes.php?rfc=rfc2474&lang=es)

¹⁹ FELICI, Santiago. Evaluación de Mecanismos de Calidad de Servicio en los routers para servicios multimedia. [Http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt](http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt)

Gráfica 2
Modelo IntServ / RSVP



RSVP es un protocolo de señalización de la capa de transporte que se utiliza para garantizar calidad de servicio (QoS) y se encarga de gestionar reservas de recursos a lo largo de todos los nodos de la red bajo la arquitectura de servicios integrados (IntServ). Por lo tanto debe ejecutarse en todos ellos para que la calidad se garantice.

Cuando es aceptada una transmisión por parte del módulo RSVP del emisor, se envía un paquete que informa a los nodos intermedios que se va a proceder a transmitir. Una vez que el receptor es avisado, envía al emisor una confirmación. Durante este proceso se hace una reserva de recursos en cada nodo intermedio para la comunicación. Esta reserva debe refrescarse periódicamente.

Para la realización de reservas es necesario que haya una autorización y que haya recursos suficientes en cada uno de los nodos del camino para que dicha reserva sea atendida.

Como ejemplo para realizar una sesión de videoconferencia se utiliza este mecanismo, siempre que la red lo soporte, para intentar garantizar un determinado ancho de banda. Los mensajes RSVP se mandan a lo largo de toda la sesión cada cierto tiempo para refrescar en los nodos el estado de la reserva.

“RSVP puede ser utilizado tanto por hosts como por routers para pedir o entregar niveles específicos de calidad de servicio (QoS) para los flujos de datos de las aplicaciones. RSVP define como deben hacer las reservas las aplicaciones y

como liberar los recursos reservados una vez que han terminado. Fue diseñado para inter operar con los actuales y futuros protocolos de enrutamiento²⁰.

1.1.1. Tipos de servicio en IntServ

Servicio	Características
Garantizado RFC2211	- Este servicio proporciona un nivel de ancho de banda y un límite en el retardo, garantizando la no existencia de pérdidas en colas. - Está pensado para aplicaciones con requerimientos en tiempo real, tales como ciertas aplicaciones de audio y vídeo. - Cada router asigna un ancho de banda y un espacio en buffer para un flujo específico.
Carga controlada (Controlled Load) RFC2212	- Este servicio no ofrece garantías en la entrega de los paquetes. - Es adecuado para aquellas aplicaciones que toleren una cierta cantidad de pérdidas y un retardo mantenidos en un nivel razonable. - Los routers que implementen este servicio deben verificar que el tráfico recibido siga las especificaciones dadas por el Tspec (especifica el trafico a transmitir), y cualquier tráfico que no las cumpla será reenviado por la red, como tráfico best-effort.
Best Effort	• Ninguna garantía (como antes sin QoS)

²⁰ FELICI, Op Cit. Pág: 23-27.

1.1.2. Características del Protocolo RSVP

1. RSVP no es un protocolo de enrutamiento, pero trabaja protocolos de enrutamiento actuales y futuros.
2. RSVP está orientado hacia el receptor: es el receptor de un flujo de datos el que inicia y mantiene la reserva de recursos para ese flujo.
3. RSVP es *soft state* (la reserva en cada nodo necesita refresco periódico), mantiene solo temporalmente el estado de las reservas de recursos del host y de los routers, de aquí que soporte cambios dinámicos de la red.
4. RSVP proporciona varios estilos de reserva (un conjunto de opciones) y permite que se añadan futuros estilos al protocolo para permitirle adaptarse a diversas aplicaciones.

1.1.3. Componentes del protocolo RSVP

Para implementar RSVP los routers han de incorporar cuatro elementos:

- *Admission Control*: comprueba si la red tiene los recursos suficientes para satisfacer la petición. Equivalente al CAC (Connection Admission Control) de ATM.
- *Policy Control*: determina si el usuario tiene los permisos adecuados para la petición realizada (por ejemplo si tiene crédito disponible). La comprobación se puede realizar consultando una base de datos mediante el protocolo COPS (Common Open Policy Service).
- *Packet Classifier*: clasifica los paquetes en categorías de acuerdo con la QoS a la que pertenecen. Cada categoría tendrá una cola y un espacio propio para buffers en el router.
- *Packet Scheduler*: organiza el envío de los paquetes dentro de cada categoría (cada cola).

Por tanto cada router ha de mantener el detalle de todas las conexiones activas que pasan por él, y los recursos que cada una ha reservado. El router mantiene información de estado sobre cada flujo que pasa por él. Además si no se pueden asegurar las condiciones pedidas se rechaza la llamada (control de admisión).

1.1.4. Desventajas del protocolo RSVP

Este protocolo presenta problemas de escalabilidad debido a la necesidad de mantener información de estado en cada router. Esto hace inviable usar RSVP en grandes redes, por ejemplo en el 'núcleo' de Internet.

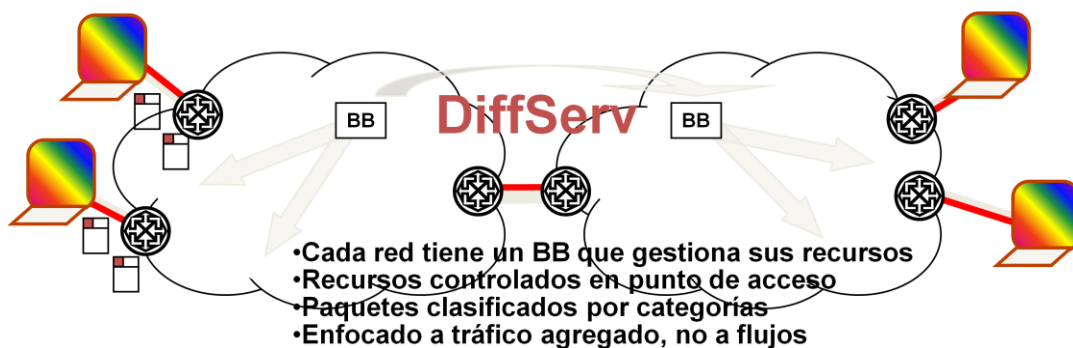
Los fabricantes de routers no han desarrollado implementaciones eficientes de RSVP, debido al elevado costo que tiene implementar en hardware las funciones de mantenimiento de la información de estado.

A pesar de todo RSVP/IntServ puede desempeñar un papel en la red de acceso, donde los enlaces son de baja capacidad y los routers soportan pocos flujos.

Recientemente ha resurgido el interés por RSVP por su aplicación en MPLS y funciones de ingeniería de tráfico. En estos casos el número de flujos no suele ser muy grande.

1.2. Modelo DiffServ (Differentiated Services – Diferenciación de servicios)²¹

Gráfica 3
Modelo DiffServ



Tomada de: <http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt>

²¹ BLAKE S, BLACK D, CARLSON M. An Architecture for Differentiated Services. 1998. <http://tools.ietf.org/html/rfc2475>

Este modelo presenta las siguientes características:

- Intenta evitar los problemas de escalabilidad que plantea IntServ/RSVP.
- Se basa en el marcado de paquetes únicamente. No hay reserva de recursos por flujo, no hay protocolo de señalización, no hay información de estado en los routers.
- Las garantías de calidad de servicio no son tan severas como en IntServ pero en muchos casos se consideran suficientes.
- En vez de distinguir flujos individuales clasifica los paquetes en categorías (según el tipo de servicio solicitado).
- A cada categoría le corresponde un SLA (Service Level Agreement). Los usuarios pueden contratar o solicitar un determinado caudal en la categoría que deseen con el operador de la red.
- Los routers tratan cada paquete según su categoría (que viene marcada en la cabecera del paquete). El Policy Control/Admission Control sólo se ha de efectuar en los routers de entrada a la red del proveedor y en los que atraviesan fronteras entre proveedores diferentes (normalmente en las fronteras entre sistemas autónomos).
- La información se puede resumir fácilmente ya que todos los flujos quedan clasificados en alguna de las categorías existentes.
- El número de categorías posibles es limitado e independiente del número de flujos o usuarios; por tanto la complejidad es constante, no proporcional al número de usuarios (decimos que la arquitectura es 'escalable').
- La información de QoS no está en los routers sino que se encuentra en los datagramas.

Gráfica 4
Cabecera Ipv4 antes de DiffServ

Version	Lon.Cab.	TOS	Longitud total			
Identificación			X	D	M	Desplazamiento fragmento
				F	F	
Tiempo de vida		Protocolo	Checksum			
Dirección de origen						
Dirección de destino						
Opciones						

Tomada de: [http:// informatica.uv.es/doctorado/SST/docto-2-qos.ppt](http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt)

Gráfica 5
Cabecera Ipv4 con DiffServ

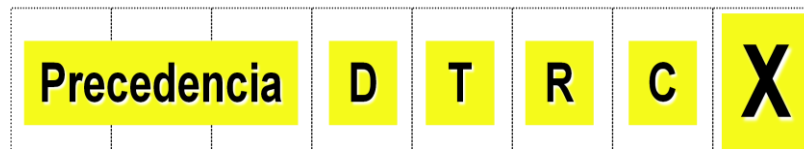
Version	Lon.Cab.	DS	Longitud total			
Identificación			X	D	M	Desplazamiento fragmento
				F	F	
Tiempo de vida		Protocolo	Checksum			
Dirección de origen						
Dirección de destino						
Opciones						

Tomada de: [http:// informatica.uv.es/doctorado/SST/docto-2-qos.ppt](http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt)

El campo TOS (obsoleto) está compuesto por los siguientes subcampos:

Gráfica 6

Campo TOS



Tomada de: [http:// informatica.uv.es/doctorado/SST/docto-2-qos.ppt](http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt)

- Presedencia: prioridad (Ocho niveles)
- D,T,R,C: flags para indicar las rutas que se quieren utilizar:
 - D: Delay (Mínimo Retardo)
 - T: Throughput (Máximo rendimiento)
 - R: Reliability (Máxima fiabilidad)
 - C: Cost (Mínimo Costo)
- X: Bit reservado.

Dicho campo pasa a ser remplazado por el campo DS (RFC 2474) que está compuesto por:

Gráfica 7

Campo DS



Tomada de: [http:// informatica.uv.es/doctorado/SST/docto-2-qos.ppt](http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt)

DSCP: Differentiated Services CodePoint. Seis bits que indican el tratamiento que debe recibir este paquete en los routers.

CU: Currently Unused (reservado). Este campo se utiliza actualmente para control de congestión.

La implementación de DiffServ en los routers se puede apreciar en la gráfica 8.

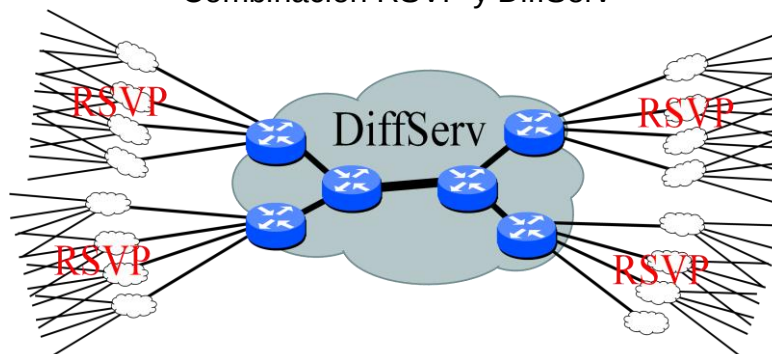
Gráfica 8
Implementación de DiffServ en Routers



Tomada de: <http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt>

1.3. Combinación de RSVP y DiffServ²²

Gráfica 9
Combinación RSVP y DiffServ



Tomada de: <http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt>

²² FELICI, Op Cit. Pág:12.

En la periferia de la red el uso de RSVP no plantea problemas y puede ser necesaria la reserva estricta de recursos. Mientras el router que conecta con el núcleo de la red IP, traduce la petición al servicio Diffserv más apropiado.

Existen varios estudios acerca de este protocolo donde se proponen varias mejoras, entre ellas sobresalen:

- Nuevas extensiones para que existan mensajes de señalización que permitan que los servidores ajusten o modifiquen las reservas, haciendo más flexible el uso del ancho de banda del enlace.
- Definir una extensión novel llamada “On Board RSVP Protocol” que soporte eficientemente comunicación end to end reservando recursos para la movilidad de los routers.
- Nuevas extensiones que permitan servicios en tiempo real en ambientes jerárquicos Mobile IPV6.

2. MANEJO DE ETIQUETAS EN LAS REDES IP – (Protocol LDP)²³

Este protocolo es utilizado en la tecnología **MPLS** (siglas de *Multiprotocol Label Switching*), conmutación de etiquetas multiprotocolo, mecanismo de transporte de datos estándar creado por la IETF (Internet Engineering Task Force) (en español Grupo de Trabajo en Ingeniería de Internet)) el cual permite transportar diferentes tipos de tráfico, como voz y paquetes IP.

Este protocolo permite asignar etiquetas a cada uno de los circuitos virtuales de origen a destino (Label Switching Path o LSP) con recursos reservados y parámetros garantizados, como parte de la información que se maneja en cada etiqueta encontramos:

- Prefijos de dirección de los routers de origen y destino.
- Versión del protocolo usado.
- Longitud del mensaje.
- Referenciación del tipo de mensaje.
- Información sobre estado de la sesión.

Esto permite mantener un canal virtual mientras se está realizando la comunicación, en el momento que termina dicha comunicación el canal es liberado permitiendo optimizar recursos ya que se elimina la subutilización cuando no se está llevando a cabo transmisión de datos.

Existen varios estudios que han permitido perfeccionar el desempeño del protocolo LDP, el primero permite pasar de un sistema que tenía poca tolerancia a fallas a uno que logra recuperarse de los problemas que se le presentan, definiendo dentro del funcionamiento y configuración del protocolo dos (2) elementos, uno detector, que identifica la falla presentada y los cambios que se dieron en la información transportada y otro elemento corrector, que recupera la información inicial y reanuda nuevamente la sesión. El segundo permite llevar a cabo una sincronización entre la información generada en la etiqueta por el nodo principal que envió la información y cada uno de los nodos por los cuales pasa el circuito virtual de comunicación, para que en caso de que alguno de esos nodos tenga alguna falla, al reiniciar el servicio automáticamente detecte la información de la etiqueta y pueda seguir manteniendo el circuito virtual.

²³ ANDERSON L, DOOLAN P, FELDMAN L. Op. Cit. Pag: 5-10.

3. PROTOCOLOS DE ENRUTAMIENTO ACTUALES EN REDES IP

3.1. Protocolo OSPF (Open Short Path First – Abrir primero la ruta de acceso más corta)²⁴

Es un protocolo de enrutamiento interno basado en el algoritmo Short Path First, estándar de Internet, que ha sido desarrollado por un grupo de trabajo del Internet Engineering Task Force - IETF, cuya especificación viene recogida en el RFC 2328.

Al ser un protocolo de routing interno, este distribuye información entre routers que pertenecen al mismo Sistema Autónomo.

El desarrollo de OSPF por parte del IETF se basa fundamentalmente en la introducción de una algoritmia diferente de la utilizada hasta el momento en los protocolos estándar de routing interno en TCP/IP para el cálculo del camino mínimo entre dos nodos de una red. A diferencia de RIP, este protocolo no envía la cantidad de saltos que separa a los routers cercanos, sino el estado de la conexión entre ellos.

3.1.1. Tipos de mensajes del Protocolo OSPF

- Hello – Saludo: se utiliza para identificar a los vecinos y así crear una base de datos en mapa local. Para enviar señales de “Estoy Vivo” al resto de routers para mantener el mapa local. Y a su vez permite elegir y encontrar un router designado para una red multienvío.
- Database Description Packets o Descripción de la base de datos: este mensaje es utilizado con el fin que un router pueda descubrir los datos que le faltan durante la fase de inicialización o sincronización una vez que dos nodos han establecido una conectividad.
- Link State Request o Petición del estado del enlace: se usa para pedir datos que un router se ha dado cuenta que le faltan en su base de datos o que están obsoletos durante la fase de intercambio de información entre dos routers.

²⁴ MOY, J. Op. Cit. Págs: 50-90.

- Link State Request o Actualización del estado del enlace: se usa como respuesta a los mensajes de Petición de estado del enlace y también para informar dinámicamente de los cambios en la topología de la red. El emisor retransmitirá hasta que se confirme con un mensaje de ACK.
- Link State ACK o ACK del estado del enlace: se usa para confirmar la recepción de una Actualización del estado del enlace.

3.1.2. Funcionamiento básico de OSPF

El fundamento principal en el cual se basa un protocolo de estado de enlace es en la existencia de un mapa de la red el cual es poseído por todos los nodos y que regularmente es actualizado.

Para llevar a cabo este propósito la red debe de ser capaz de entre otros objetivos de:

- Almacenar en cada nodo el mapa de la red.
- Ante cualquier cambio en la estructura de la red actuar rápidamente, con seguridad si crear bucles y teniendo en cuenta posibles particiones o uniones de la red.

3.1.3. Mapa de Red Local

La creación del mapa de red local en cada router de la red se realiza a través de una tabla donde:

Fila: representa a un router de la red y cualquier cambio que le ocurra a ese router será reflejado en este registro de la tabla a través de los registros de descripción.

Columna: representa los atributos de un router que son almacenados para cada nodo. Entre los principales atributos por nodo tenemos: un identificador de interfase, el número de enlace e información acerca del estado del enlace, o sea, el destino y la distancia o métrica.

Ejemplo: Cada uno de los siguientes routers representados por letras se encuentra enlazado entre sí de acuerdo a un determinado número de enlaces.

A --- 1 --- B --- 2 --- C --- 4 --- D --- 3 --- A

Teniendo en cuenta el número de enlaces y de saltos existentes entre los routers, el Mapa de Red Local sería:

Origen	Destino	Enlace	Distancia (Saltos)
A	B	1	1
B	C	2	1
C	D	4	1
D	A	3	1
B	A	1	1
C	B	2	1
D	C	4	1
A	D	3	1

Un cambio en la topología de la red es detectado en primer lugar o por el nodo que causo el cambio o por los nodos afectados por el enlace que provoco el cambio. El protocolo o mecanismo de actualización la información por la red debe ser rápido y seguro, y estos son los objetivos del protocolo de inundación y de intercambio o sincronización empleado en OSPF. Dicho protocolo se denomina Protocolo de Inundación: The flooding Protocol.

3.1.4. Características de OSPF

- Respuesta rápida y sin bucles ante cambios.

Debido a que todos los nodos de la red calculan el mapa de manera idéntica y poseen el mismo mapa no se generan bucles ni nodos que se encuentren contando en infinito; principal problema sufrido por los protocolos basados en la algoritmia de vector distancia como RIP.

- Seguridad ante los cambios.

Para que el algoritmo de routing funcione adecuadamente debe existir una copia idéntica de la topología de la red en cada nodo de esta.

El protocolo OSPF especifica que todos los intercambios entre routers deben ser autenticados. El OSPF permite una variedad de esquemas de autenticación y también permite seleccionar un esquema para un área diferente al esquema de otra área. La idea detrás de la autenticación es garantizar que sólo los routers confiables difundan información de routing.

- Soporte de múltiples métricas.

La tecnología actual hace que sea posible soportar varias métricas en paralelo. Evaluando el camino entre dos nodos en base a diferentes métricas es tener distintos mejores caminos según la métrica utilizada en cada caso, pero surge la duda de cuál es el mejor. Esta elección se realizara en base a los requisitos que existan en la comunicación.

Diferentes métricas utilizadas pueden ser:

Mayor rendimiento
Menor retardo
Menor costo
Mayor fiabilidad

La posibilidad de utilizar varias métricas para el cálculo de una ruta, implica que OSPF provea de un mecanismo para que una vez elegida una métrica en un paquete para realizar su routing esta sea la misma siempre para ese paquete, esta característica dota a OSPF de un routing de servicio de tipo en base a la métrica.

- Balanceado de carga en múltiples caminos

OSPF permite el balanceado de carga entre los nodos que exista más de un camino. Para realizar este balanceo aplica:

Una versión de SPF con una modificación que impide la creación de bucles parciales y luego un algoritmo que permite calcular la cantidad de trafico que debe ser enviado por cada camino.

- Escalabilidad en el crecimiento de rutas externas

El continuo crecimiento de Internet es debido a que cada vez son más los sistemas autónomos que se conectan entre sí a través de routers externos.

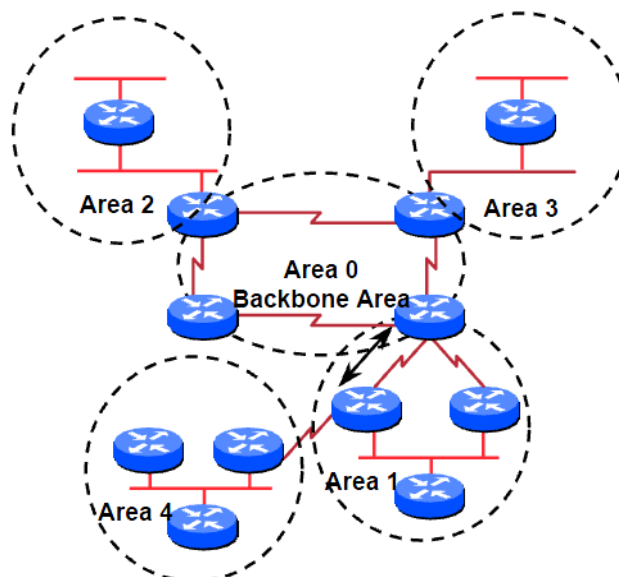
Además de tener en cuenta la posibilidad de acceder al exterior del sistema autónomo a través de un determinado router externo u otro se debe tener en cuenta que se tiene varios proveedores de servicios y es más versátil elegir en cada momento el router exterior y servicio requerido que establecer una ruta y servicio por defecto cuando se trata de routing externo como se tenía hasta ahora.

OSPF soluciona este problema permitiendo tener en la base de datos del mapa local los denominados “gateway link state records”. Estos registros nos permiten almacenar el valor de las métricas calculadas y hacen más fácil el cálculo de la ruta óptima para el exterior. Por cada entrada externa existirá una nueva entrada de tipo “gateway link state records” en la base de datos, es decir, la base de datos crecerá linealmente con el número de entradas.

- Routing Jerárquico

OSPF establece una jerarquía en la red y la parte en “áreas”, existiendo una área especial denominada “backbone área”. (Ver gráfica 10)

Gráfica 10
Áreas OSPF



Tomada de: http://ws.edu.isoc.org/data/2004/Introduccion_OSPF.pdf

En un “área” se aplica el protocolo OSPF de manera independiente como si de una red aislada se tratase, es decir, los routers del área solo contiene en su mapa local la topología del área, así que el coste en calculo es proporcional al tamaño del área y no de la totalidad de la red.

Cada área incluye un conjunto de subredes IP. La comunicación entre routers de un área se resuelve directamente a través del mapa local de área que cada router posee.

Estas áreas se conectan entre sí a través del “backbone área”, mediante routers que pertenecen normalmente a una “área” y al “backbone área”. Estos routers se denominan “area-border routers” y como mínimo existe uno entre un área y el backbone.

3.1.5. Relación con los vecinos en OSPF

Cada enrutador OSPF realiza un seguimiento de sus nodos vecinos, estableciendo distintos tipos de relación con ellos. Respecto a un enrutador dado, sus vecinos pueden encontrarse en siete estados diferentes:

- **Estado Desactivado (DOWN)**

En el estado desactivado, el proceso OSPF no ha intercambiado información con ningún vecino. OSPF se encuentra a la espera de pasar al siguiente estado (Estado de Inicialización)

- **Estado de Inicialización (INIT)**

Los *routers* (enrutadores) OSPF envían paquetes tipo 1, o paquetes Hello, a intervalos regulares con el fin de establecer una relación con los Routers vecinos. Cuando una interfaz recibe su primer paquete Hello, el *router* entra al estado de Inicialización. Esto significa que este sabe que existe un vecino a la espera de llevar la relación a la siguiente etapa.

Los dos tipos de relaciones son Bidireccional y Adyacencia. Un *router* debe recibir un paquete Hello (Hola) desde un vecino antes de establecer algún tipo de relación.

- **Estado Bidireccional (TWO-WAY)**

Empleando paquetes Hello, cada enrutador OSPF intenta establecer el estado de comunicación bidireccional (dos-vías) con cada enrutador vecino en la misma red IP. Entre otras cosas, el paquete Hello incluye una lista de los vecinos OSPF conocidos por el origen. Un enrutador ingresa al estado Bidireccional cuando se ve a sí mismo en un paquete Hello proveniente de un vecino.

El estado Bidireccional es la relación más básica que vecinos OSPF pueden tener, pero la información de encaminamiento no es compartida entre estos. Para aprender los estados de enlace de otros enrutadores y eventualmente construir una tabla de enrutamiento, cada enrutador OSPF debe formar a lo menos una adyacencia. Una adyacencia es una relación avanzada entre enrutadores OSPF que involucra una serie de estados progresivos basados no sólo en los paquetes Hello, sino también en el intercambio de otros 4 tipos de paquetes OSPF. Aquellos enrutadores intentando volverse adyacentes entre ellos intercambian información de encaminamiento incluso antes de que la adyacencia sea completamente establecida. El primer paso hacia la adyacencia es el estado ExStart.

- **Estado EXSTART**

Técnicamente, cuando un enrutador y su vecino entran al estado ExStart, su conversación es similar a aquella en el estado de Adyacencia. ExStart se establece empleando descripciones de base de datos tipo 2 (paquetes DBD), también conocidos como DDPs. Los dos enrutadores vecinos emplean paquetes Hello para negociar quien es el "maestro" y quien es el "esclavo" en su relación y emplean DBD para intercambiar bases de datos.

Aquel enrutador con el mayor router ID "gana" y se convierte en el maestro. Cuando los vecinos establecen sus roles como maestro y esclavo entran al estado de Intercambio y comienzan a enviar información de encaminamiento.

- **Estado de Intercambio (EXCHANGE)**

En el estado de intercambio, los enrutadores vecinos emplean paquetes DBD tipo 2 para enviarse entre ellos su información de estado de enlace. En otras palabras,

los enrutadores se describen sus bases de datos de estado de enlace entre ellos. Los enrutadores comparan lo que han aprendido con lo que ya tenían en su base de datos de estado de enlace. Si alguno de los enrutadores recibe información acerca de un enlace que no se encuentra en su base de datos, este envía una solicitud de actualización completa a su vecino. La información completa de encaminamiento es intercambiada en el estado cargando.

- **Estado Cargando (LOADING)**

Después de que las bases de datos han sido completamente descritas entre vecinos, estos pueden requerir información más completa empleando paquetes tipo 3, requerimientos de estado de enlace (LSR). Cuando un enrutador recibe un LSR este responde empleando un paquete de actualización de estado de enlace tipo 4 (LSU). Estos paquetes tipo 4 contienen las publicaciones de estado de enlace (LSA) que son el corazón de los protocolos de estado de enlace. Los LSU tipo 4 son confirmados empleando paquetes tipo 5 conocidos como confirmaciones de estado de enlace (LSAcks).

- **Estado de Adyacencia completa (FULL)**

Cuando el estado de la carga ha sido completado, los enrutadores se vuelven completamente adyacentes. Cada enrutador mantiene una lista de vecinos adyacentes, llamada base de datos de adyacencia.

3.2. Border Gateway Protocol (BGP) – Protocolo de Gateway Fronterizo²⁵

El Border Gateway Protocol (Protocolo de Gateway fronterizo) es un protocolo de enrutamiento por vector de distancia usado para enrutar paquetes entre dominios, gestiona el enrutamiento entre 2 o más routers fronterizos de sistemas autónomos, administra eficientemente la agregación y propagación de rutas entre dominios.

²⁵ REKHTER, Y. Op Cit. Págs: 3-18.

Mediante este protocolo se intercambian prefijos entre los ISP (proveedores de servicio de internet) existentes.

BGP se diseñó para permitir la cooperación en el intercambio de información de enrutamiento entre dispositivos de encaminamiento, llamados pasarelas, en sistemas autónomos diferentes. El protocolo opera en términos de mensajes, que se envían utilizando TCP. El repertorio de mensajes es el siguiente:

1. OPEN (aparece cuando se abre una conexión TCP con el vecino)
2. UPDATE (proporciona una lista de rutas a usar por el dispositivo)
3. KEEPALIVE (lo envía regularmente para indicar actividad)
4. NOTIFICACION (se envía al presentarse algún tipo de error)

BGP supone tres procedimientos funcionales:

- Adquisición de vecino.
- Detección de vecino alcanzable.
- Detección de red alcanzable.

Dos dispositivos de encaminamiento se considera que son vecinos si están en la misma subred. Si los dos dispositivos de encaminamiento están en sistemas autónomos, podrían desear intercambiar información de encaminamiento. Para esto es necesario realizar primero el proceso de adquisición de vecino. Se requiere un mecanismo formal de encaminamiento ya que alguno de los dos vecinos podría no querer participar. Existirán situaciones en las que un vecino no desee intercambiar información esto se puede deber a múltiples factores como por ejemplo que este sobresaturado y entonces no quiere ser responsable del tráfico que llega desde fuera del sistema.

3.3. Routing Information Protocol RIP – Protocolo de encaminamiento de información²⁶

Es un protocolo de encaminamiento interno, funciona en la parte interna de la red, la cual no está conectada al backbone de internet, normalmente usado en los routers, además calcula el camino más corto hacia la red de destino usando el algoritmo del vector de distancias, el cual tiene en cuenta la dirección y la distancia hacia cualquier enlace de la red. La distancia o métrica está determinada por el número de saltos, que indica el número de routers por los cuales tiene que pasar hasta alcanzar la red de destino.

²⁶ MALKIN, G. Op. Cit. Págs: 3-5.

Este protocolo es fácil de configurar. Además, es un protocolo abierto, soportado por muchos fabricantes pero tiene la desventaja que, para determinar la mejor métrica, únicamente toma en cuenta el número de saltos (por cuántos routers o equipos similares pasa la información); no toma en cuenta otros criterios importantes, como por ejemplo ancho de banda de los enlaces. Por ejemplo, si tenemos una métrica de 2 saltos hasta el destino con un enlace de 64 Kbps y una métrica de 3 saltos, pero con un enlace de 2 Mbps, lamentablemente RIP tomara el enlace de menor número de saltos aunque sea el más lento. Protocolo usado pero con limitaciones, limitando máximo a 15 saltos.

RIP genera un alto volumen de tráfico al enviar toda la tabla de ruteo en cada actualización, además utiliza métricas fijas para comparar rutas alternativas, lo que lo limita para desempeñarse en escenarios con parámetros de tiempo real donde sea necesario tener en cuenta retardos o cargas en los enlaces.

La tabla de ruteo que se encuentra en cada uno de los hosts de la red y que hacen uso del protocolo RIP, tiene los siguientes campos:

- Dirección de destino (no permite subneting en la versión 1)
- Siguiendo salto (siguiendo enrutador vecino al que pasara)
- Interfaz de salida del enrutador.
- Métrica (Conteo de saltos).
- Temporizador (desde la última actualización de rutas).

El RIP es un protocolo Vector distancia que debido al crecimiento de las redes no es escalable.

3.4. Intermediate System to Intermediate System (IS-IS)²⁷

En los últimos años el protocolo de enrutamiento de Sistema Intermedio a Sistema Intermedio (IS-IS) ha comenzado a incrementar su popularidad. IS-IS converge rápidamente y es muy escalable. Es un protocolo OSI de encaminamiento jerárquico de pasarela interior o IGP (Interior Gateway Protocol), que usa el estado de enlace para encontrar el camino más corto mediante el algoritmo SPF (Shortest Path First).

²⁷ ORAN D. OSI IS-IS Intra-domain Routing Protocol. 1990.
[Http://tools.ietf.org/html/rfc1142](http://tools.ietf.org/html/rfc1142).

Características de IS-IS

- Enrutamiento jerárquico
- Comportamiento sin clases
- Inundación rápida de nueva información
- Convergencia rápida
- Muy escalable
- Sintonizador de tiempo flexible

Características de los servicios de red OSI

- Independencia de la infraestructura de comunicación de capas inferiores
- Transferencia end-to-end
- Transparencia
- Selección de calidad de servicio (QoS)
- Direccionamiento

El conjunto de protocolos OSI soporta numerosos protocolos estándares en la capa física, enlace, red, transporte, sesión, presentación y aplicación.

Especificado originalmente por la ISO, es un protocolo dinámico, de estado y enlace, ínter dominio y de interior (IGP). Opera en un ambiente de servicio sin conexión OSI (CLNS), seleccionando rutas sobre una métrica de costo asignada a los enlaces por un administrador como el valor de la ruta a un router vecino.

El gran dominio puede ser dividido en áreas. Cada sistema reside en un área. El enrutamiento dentro del área es referido como enrutamiento Nivel 1, y entre áreas como Nivel 2. Un sistema intermedio (IS) nivel 2 mantiene la pista de las rutas a los destinos de las otras áreas. Un nivel 1 mantiene la pista del enrutamiento dentro del área. Cuando un paquete lleva como destino otra área, el nivel 1 envía el paquete al nivel 2 más cercano dentro de su área. Sin importar el área de destino, donde él puede viajar por caminos de enrutamiento nivel 1 hasta el destino.

También es utilizado con el protocolo TCP/IP. De tal manera que es capaz de encaminar paquetes IP y CLNP (ConnectionLess Network Protocol). CLNP es un protocolo de capa de red OSI que transporta datos de capa superior e indicaciones de errores sobre enlaces sin conexión. CLNP provee la interface entre CLNS y las capas superiores. CLNS OSI es un servicio de capa de red similar al servicio IP crudo que se comunica sobre protocolos de red sin conexión CLNP.

CLNS confía en los protocolos de la capa de transporte para ejecutar detección y corrección de errores.

No emplea encapsulación para los paquetes, ni ninguna diferencia relevante entre ellos, excepto que en IP añade información adicional.

El protocolo tiene un gran parecido con OSPF ya que en ambos se utiliza el estado de enlace para la búsqueda de caminos (utilizan puentes designados para eliminar bucles) y la asignación de redes en grupos para mejorar la eficiencia de la red. Pero IS-IS tiene ciertas ventajas respecto a OSPF tales como compatibilidad con IPv6 o que permite conectar redes con protocolos de encaminamiento distintos.

4. MULTIPROTOCOL LABEL SWITCHING MPLS – CONMUTACIÓN DE ETIQUETAS MULTIPROTOCOLO²⁸

La esencia de MPLS (***Multiprotocol Label Switching***) es la generación de pequeñas etiquetas de longitud fija que actúan como una representación de la cabecera del paquete IP. Los paquetes IP tienen un campo en la cabecera que contiene la dirección a la que el paquete debe ser enviado. Las redes de encaminamiento tradicionales procesan esta información en cada uno de los routers que atraviesa el paquete (encaminamiento hop-by-hop). En MPLS, los paquetes IP son encapsulados con estas etiquetas por el primer dispositivo MPLS que encuentre al entrar en la red. Este dispositivo, que recibe el nombre de LER (***Label Edge Router***), analiza el contenido de la cabecera IP y selecciona una etiqueta apropiada con la cual encapsular el paquete. Una parte del gran poder de MPLS se debe al hecho de que, en contraste con el encaminamiento IP tradicional, este análisis puede ser basado en algo más que la dirección de destino que hay en la cabecera IP. En todos los nodos subsiguientes de la red, la etiqueta MPLS (y no la cabecera IP) es usada para decidir por donde se envía el paquete. Finalmente, cuando los paquetes MPLS abandonan la red, otro LER elimina las etiquetas.

Cuando se quiere llevar a cabo la transmisión de información en una red IP, esta transmisión debe realizarse sobre una tecnología de segundo nivel (ATM, Frame Relay, PPP, Ethernet, etc.) y MPLS actúa como integración entre el tercer nivel IP con el segundo nivel.

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 clases de servicio, 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 archivos (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 (Type of Service), rebautizado en DiffServ como el octeto DS.

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

²⁸ ROSEN E, VISWANATHAN A. Op. Cit. Pag: 5-34.

1. 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.
2. 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. Ejemplo: un LSP puede ser para tráfico de máxima prioridad, otro para una prioridad media y un tercero para tráfico best-effort, tres niveles de servicio que lógicamente tendrán distintos precios.

Por lo tanto, MPLS sirve para la administración de la calidad de servicio al definir 5 clases de servicios, conocidos como CoS.

1. *Video*: La clase de servicio para transportar video tiene un nivel de prioridad más alto que las clases de servicio para datos.

2. *Voz*: La clase de servicio para transportar voz tiene un nivel de prioridad equivalente al de video, es decir, más alto que las clases de servicio para datos.

3. *Datos de alta prioridad (D1)*: Ésta es la clase de servicio con el nivel de prioridad más alto para datos. Se utiliza particularmente para aplicaciones que son críticas en cuanto necesidad de rendimiento, disponibilidad y ancho de banda.

4. *Datos de prioridad (D2)*: Esta clase de servicio se relaciona con aplicaciones que no son críticas y que tienen requisitos particulares en cuanto a ancho de banda.

5. *Datos no prioritarios (D3)*: representan la clase de servicio de prioridad más baja.

Las especificaciones de MPLS operan en la capa 2 del modelo OSI y pueden funcionar particularmente en redes IP, ATM o frame relay.

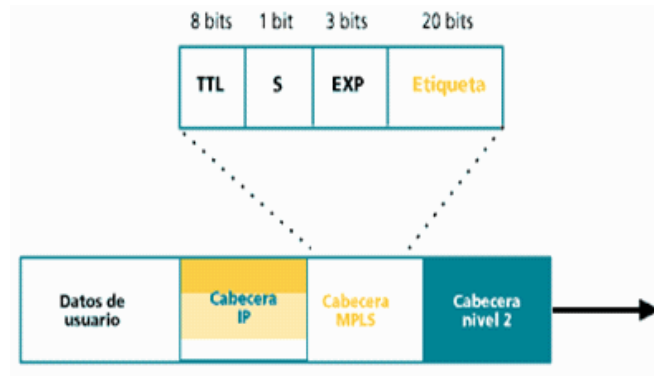
MPLS es hoy día una solución clásica y estándar al transporte de información en las redes. Aceptado por toda la comunidad de Internet, ha sido hasta hoy una solución aceptable para el envío de información, utilizando Routing de paquetes con ciertas garantías de entrega.

4.1. Características de las redes MPLS

- Se basa en la asignación e intercambio de etiquetas.
- Separa el routing del forwarding (enrutamiento / transmisión).
- Permite garantizar QoS sobre IP.
- Permite implementar ingeniería de tráfico.
- Ofrece niveles de rendimiento diferenciados y priorización del tráfico, así como aplicaciones de voz y multimedia. Y todo ello en una única red.
- Asigna a los datagramas de cada flujo una etiqueta única que permite una conmutación rápida en los routers intermedios (solo se mira la etiqueta, no la dirección de destino).
- MPLS se basa en el etiquetado de los paquetes en base a criterios de prioridad (según la clase de servicio) y/o calidad (QoS).
- MPLS es una tecnología que permite ofrecer enrutamiento con calidad de servicio (QoS), independientemente de la red sobre la que se implemente.
- El etiquetado en capa 2 permite ofrecer servicio multiprotocolo y ser portable sobre multitud de tecnologías de capa de enlace: ATM, Frame Relay, líneas dedicadas, LANs.
- La etiqueta MPLS se coloca delante del paquete de red y detrás de la cabecera de nivel de enlace. (Ver gráfica 11)

Gráfica 11

Etiqueta MPLS



Tomada de <http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt>

- El formato de la etiqueta MPLS posee 32 bits. (Ver Gráfica 12)

Gráfica 12

Formato de etiqueta MPLS



Etiqueta: La etiqueta propiamente dicha que identifica una FEC (con significado local)

Exp: Bits para uso experimental; una propuesta es transmitir en ellos información de DiffServ

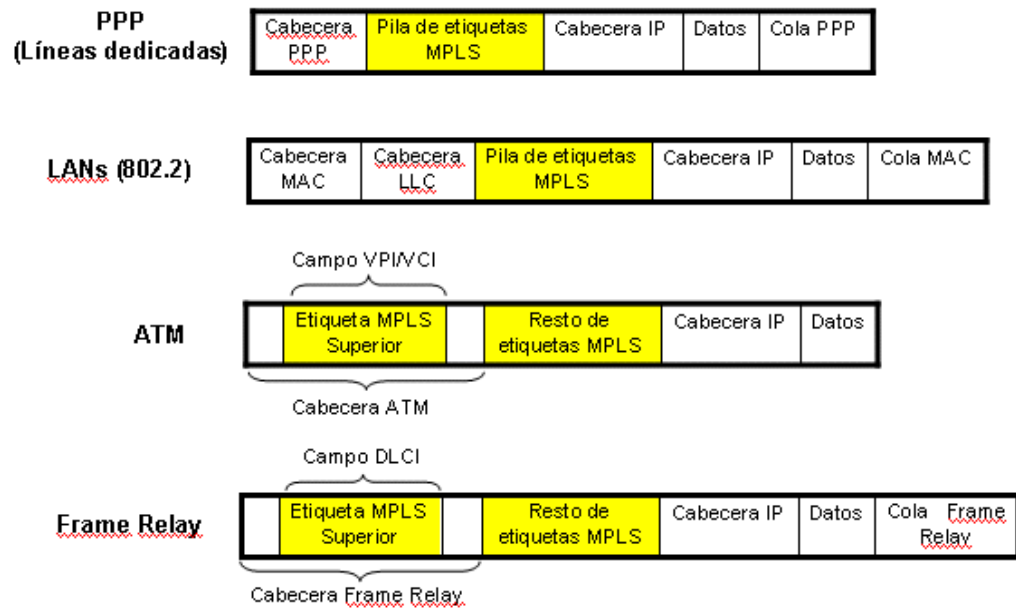
S: Vale 1 para la primera entrada en la pila (la más antigua), cero para el resto. Esta es la primera etiqueta introducida.

TTL: Contador del número de saltos. Este campo reemplaza al TTL de la cabecera IP durante el viaje del datagrama por la red MPLS.

Tomada de: informatica.uv.es/doctorado/SST/docto-2-qos.ppt

La ubicación de la etiqueta MPLS según tecnología capa 2 de enlace se puede apreciar en la siguiente gráfica:

Gráfica 13
Ubicación de etiqueta según tecnología

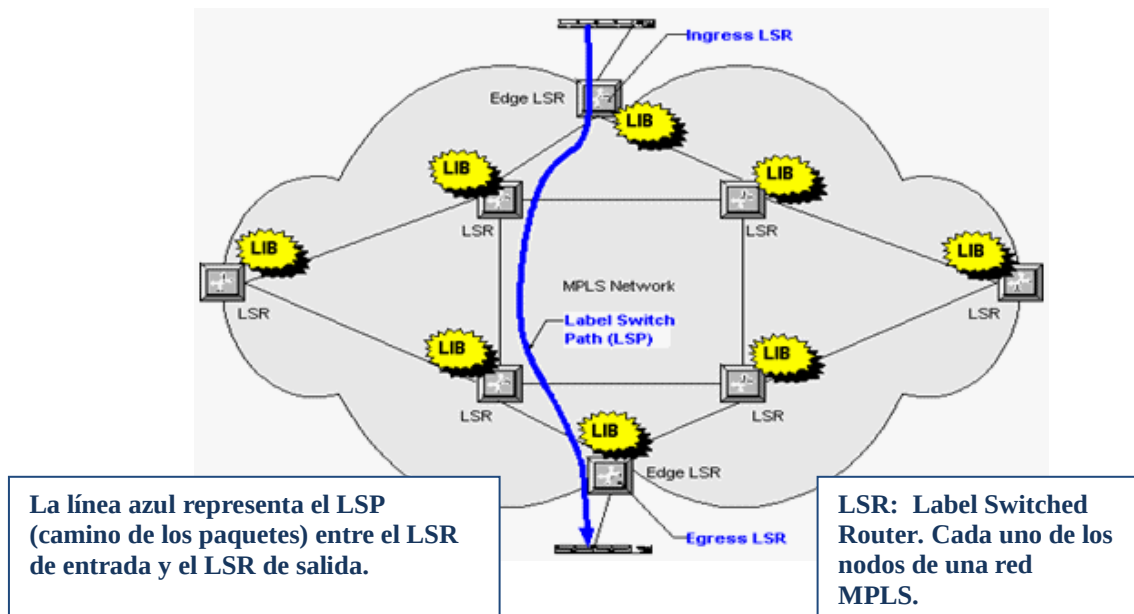


Tomada de: <http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt>

4.2. Routing en MPLS

Gráfica 14

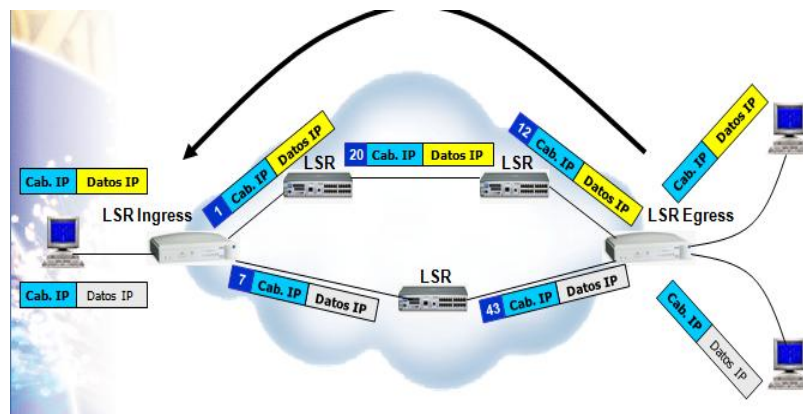
Enrutamiento en MPLS



Tomada de: <http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt>

Gráfica 15

Intercambio de Etiquetas en MPLS



Tomada de: www.eie.fceia.unr.edu.ar/~comunica/TBAA2003pub/MPLS.ppt

El enrutamiento a través de una red MPLS se hace teniendo en cuenta los siguientes criterios (Ver gráficas 14 y 15):

- Los paquetes se envían en función de las etiquetas.
- No se examina la cabecera de red completa.
- El direccionamiento es más rápido.
- Cada paquete es clasificado en unas clases de tráfico denominadas FEC (*Forwarding Equivalence Class*) grupo de paquetes IP que son retransmitidos sobre el mismo LSP con el mismo tratamiento.
- Los LSPs (*camino que siguen los paquetes*) por tanto definen las asociaciones FEC-etiqueta.

4.3. Funcionamiento de una red MPLS

El funcionamiento de la red MPLS (Ver Gráfica 15) se lleva a cabo de la siguiente manera:

- LER: recibe el tráfico a la entrada del dominio MPLS.
- Clasifica el tráfico (origen, destino, tipo de tráfico).
- Asigna una FEC (*Forwarding Equivalence Class*) grupo de paquetes IP que son retransmitidos sobre el mismo LSP con el mismo tratamiento.
- Envía el tráfico por el LSP asociado a la FEC.
- Dentro del LSP cada LSR se limita a intercambiar etiquetas.
- Si un LSR no sabe qué hacer descarta el paquete.

5. PROTOCOLOS COMPLEMENTARIOS PARA HACER INGENIERÍA DE TRÁFICO EN REDES MPLS²⁹

En este capítulo, se presentan nuevos medios que son importantes para el enrutamiento y la ingeniería de tráfico. En un ambiente IP, la ingeniería de tráfico implica determinar los pesos de los enlaces en un ambiente OSPF /IS-IS pudiéndose controlar el flujo global de la red. Sin embargo para la ingeniería de tráfico en IP todos los paquetes hacia un destino siguen el mismo camino debido a que el destino se encuentra orientado bajo el paradigma del enrutamiento IP; Ciertamente, si se encuentran caminos de igual costo, el tráfico puede tomar dos o más

caminos diferentes hacia el destino. Sin embargo teniendo un mecanismo de control bien definido el cual permita que los paquetes tomen diferentes caminos para un destino en particular, por ejemplo, dependiendo del tipo de paquete o clase de tráfico, sería lo mejor.

Esto puede ser analizado como ingeniería de tráfico con más controles para distinguirla de la ingeniería de tráfico IP basada en la configuración de pesos de enlace como control primario.

Para alcanzar la habilidad de definir un camino y la fuerza de paquetes para un flujo ó clase de tráfico en particular a transmitirse por él, se requiere un mecanismo que permita definirlo o identificarlo, independientemente de los paquetes designados para que sean transmitidos por él.

La pregunta entonces es ¿Cómo definir tales caminos? Cabe señalar que la idea de definir un camino surge de la red telefónica donde la configuración de la llamada se realiza primero usando la señalización ISUP. Para las redes IP el interés radica en tener un mecanismo que pueda trabajar con tráfico de paquetes. Es importante tener en cuenta que si se requiere un mecanismo sobre una llamada de voz que pueda también trabajar sobre un ambiente IP, una sesión del protocolo SIP puede ser considerada. SIP es extremo a extremo. Aunque el interés es de un mecanismo que sea de router a router. Por lo tanto, debe quedar claro que la creación de un camino entre dos routers, requiere de un mecanismo de señalización, el cual sin duda es necesario para configurarlo. Segundo, si parte del objetivo es hacer ingeniería de tráfico con

²⁹ MEDHI Deeppankar, RAMASAMY Karthikeyan. Network Routing: Algorithms, Protocols and Architectures. Estados Unidos de América. Morgan Kaufmann Publications. 2007. Págs: 613-626.

más controles, todavía se requiere que el estado de la red sea comunicado entre los routers.

Una forma de comunicar el estado de la red es informando sobre el estado del enlace a través de un protocolo de enrutamiento de estado del enlace. Sin embargo, los protocolos de estado de enlace estándar definido para redes IP transportan solo un valor de métrica único para cada enlace, normalmente representado por el costo del enlace. Para ingeniería de tráfico con más controles, información adicional, tales como el ancho de banda de un enlace debe ser comunicada. Esto significa que para aprender sobre el estado de una red, el paradigma del protocolo del estado del enlace podría ser usado, utilizando además información adicional acerca de cada enlace. Finalmente para diseñar una red, debería ser posible invocar diferentes sistemas que permitan calcular caminos de enrutamiento, el cual sea preferiblemente separado del mecanismo de actualización del enlace. Así un sistema de comunicación de caminos también puede depender del uso de operaciones específicas y de los requerimientos de la red para proveer sus servicios.

Con los antecedentes expuestos, podemos ahora analizar los ambientes favorables que permitan la señalización, configuración y la ingeniería de tráfico con más controles requeridos.

5.1. Extensión de Ingeniería de Tráfico para Protocolos de enrutamiento.

Se debe tener en cuenta que como mínimo, el ancho de banda de un enlace se debe tener en cuenta para efectos de aplicar la ingeniería de tráfico. Además, un enlace puede permitir mayor ancho de banda reservado debido al aumento anunciado de multiplexado estadístico para determinados tipos de tráfico, lo que significa que el número de ofertas puede ser tolerable. También, un enlace podrá tener un ancho de banda en la actualidad sin reservas, el cual es útil para los cálculos de los caminos de enrutamiento, pero que no necesariamente se basa en un cálculo de la ruta más corta. Dado que una red puede proporcionar más de un tipo de servicios priorizados, sería útil anunciar el ancho de banda sin reserva permitido para cada clase de prioridad. Además, un proveedor de red puede utilizar una métrica diferente, distinta a la relación métrica estándar; ésta métrica del enlace podría tener un significado interno es decir solamente para el proveedor.

En resumen se debe considerar para un enlace:

- Máximo ancho de banda del enlace que se puede utilizar.
- Máxima reserva de ancho de banda en caso de permitirse múltiples demandas.
- Ancho de banda sin ser reservado disponible en diferentes niveles de prioridad.
- Métricas de ingeniería de tráfico.

La pregunta es: Cómo es comunicada esta información?

Esto da lugar a la aparición de dos protocolos de enrutamiento de estados de enlace denominados OSPF / IS-IS. Estos dos protocolos se han ampliado para tener en cuenta las consideraciones anteriores para un determinado enlace.

5.2. Paquetes etiquetados y LSP.

Una etiqueta que identifica una FEC en MPLS es de 20 bits de largo y forma parte de un encabezado Shim MPLS denominado paquete MPLS. Si un paquete MPLS es para transportar un paquete IP, entonces se puede pensar que MPLS se ha introducido en la capa 2 (Enlace). La principal diferencia de MPLS de estar en la capa 2 es que proporciona una forma de enrutamiento a través de etiquetas y de LSP (camino de los paquetes) a través de múltiples saltos. Para comprender como funciona una red MPLS retomar el Capítulo IV.

El encabezado Shim Header de MPLS incluye 3 bits experimentales, un “S” bit para indicar si la etiqueta es la última en caso de existir una cola de etiquetas y 8 bits para el tiempo de vida de la etiqueta (contador de número de saltos). Ver Gráfica 12.

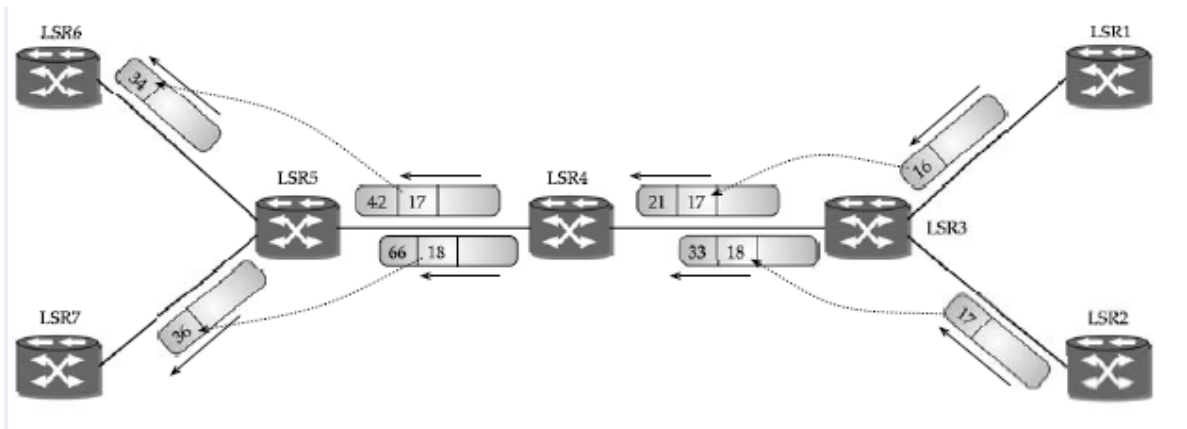
Los bits experimentales están destinados a describir los servicios que requieren diferentes prioridades. Los valores de 0 a 15 son reservados. Por ejemplo: la etiqueta con valor 0 (etiqueta nula explícita) se refiere a un paquete que se enviará sobre la base de la cabecera IPv4. Del mismo modo, la etiqueta 2 sirve de etiqueta nula explícita para IPv6. La etiqueta nula explícita puede ser utilizada por un router de salida para señalar su (penúltimo router).

Existe otro término (túnel), estrechamente relacionado con un LSP. Un túnel ofrece un servicio de transporte entre dos routers intermedios para que los paquetes de un flujo específico puedan transportarse sin intercambiar etiquetas

entre dichos routers. Un túnel en MPLS puede ser realizado por un LSP bien definido, basado en servicios con requerimientos de ingeniería de tráfico. (Ver Gráfica 16)

Gráfica 16

Caminos de etiquetas conmutadas usando un túnel LSP

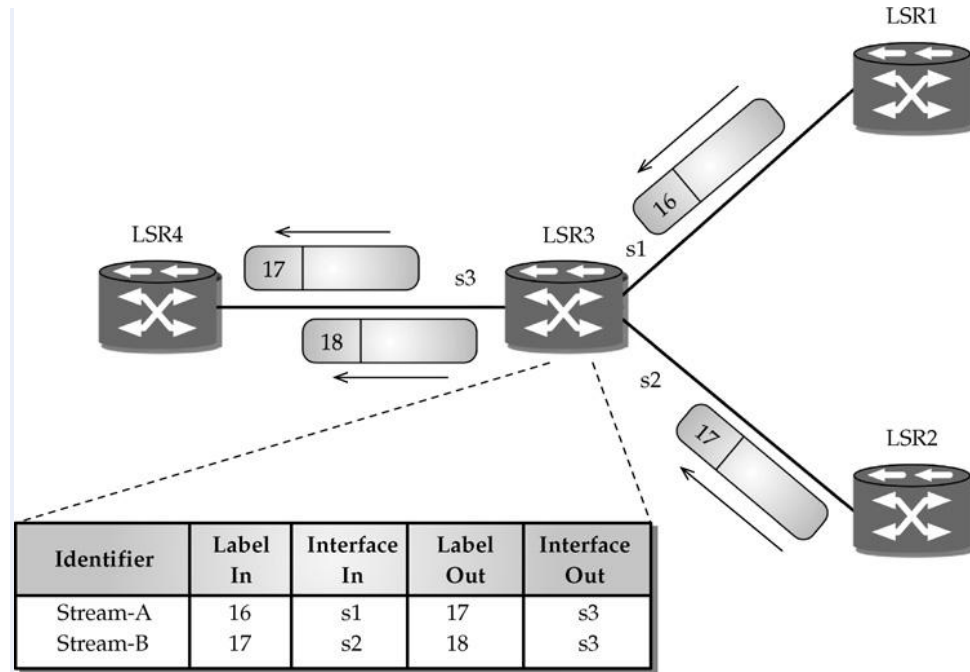


Tomada de: Network Routing: Algorithms, Protocols and Architectures. Pág. 651

Cuando un paquete MPLS llega a un LSR, la etiqueta de entrada se intercambia por una etiqueta de salida; Esto supone que un LSP ya está definido y las tablas de búsqueda en los LSR's tienen las entradas apropiadas. Luego antes de enviar un paquete recibido MPLS, el valor de TTL se reduce a uno. Si el valor de TTL se convierte en cero, entonces el paquete MPLS es abandonado. Dado que el campo TTL es de 8 bits de longitud, la probabilidad de una ruta con más de 255 saltos es cero. (Ver gráfica 17). En la gráfica se observa como un paquete con la etiqueta MPLS 16 llega a LSR3 de LSR1 y se intercambia con la etiqueta 17 para su transmisión a LSR4; del mismo modo, el paquete MPLS con la etiqueta 17 que llega desde LSR2 a LSR3 se intercambia con la etiqueta 18 para su transmisión a LSR4. En este caso, se supone que los dos LSPs, LSR1-LSR3-LSR4 y LSR2-LSR3-LSR4, ya están definidos.

Gráfica 17

Intercambio y conmutación de etiquetas entre caminos



Tomada de: Network Routing: Algorithms, Protocols and Architectures. Pág. 650

Cabe señalar que MPLS también permite apilar (almacenar) las etiquetas. Dicha acumulación de etiquetas que forman una cabecera Shim MPLS pueden aparecer más de una vez; cada una de ellas está relacionada con un LSP en particular entre determinados puntos. (Ver Gráfica 16).

Otro concepto que es asociado al LSP es el de la FEC. Cada paquete es clasificado en unas clases de tráfico denominadas FEC (*Forwarding Equivalence Class*) grupo de paquetes IP que son retransmitidos sobre el mismo LSP con el mismo tratamiento.

5.3. Distribución de Etiquetas.

Entre la definición de una FEC y el establecimiento de un LSP para esta FEC, hay otra fase muy importante conocido como la distribución de la etiqueta. La idea básica de la distribución de la etiqueta es análoga al intercambio de información en una ruta. Algo similar a BGP.

Una diferencia clave es que mientras que un router BGP calcula y determina una decisión de enrutamiento de salida, en MPLS, un LSR fundamentalmente confía en que la información de la etiqueta se basa en un LSP válido.

Tenga en cuenta que cada uno de los LSP puede ser generado por una entidad externa, pudiendo ser un director de ingeniería de tráfico. De modo que las etiquetas de distribución pueden estar asociados con LSRs, cada LSR debe tener un identificador (LSR ID) que debe ser único en el dominio MPLS. Normalmente, la dirección de un router IP sirve como ID LSR.

En MPLS, la etiqueta vinculante se refiere a asociar directamente a una etiqueta específica un determinado FEC. Consideremos dos LSRs, LSR-U y LSR-D, que se han puesto de acuerdo para obligar a una etiqueta específica para los paquetes que serán enviados de LSR-U para LSR-D; en lo que respecta a esta etiqueta obligatoria, LSR-U es denominado LSR (extremo superior) y el LSR-D (extremo inferior). El LSR intermedio informa al LSR superior de la unión. Podemos decir entonces que las etiquetas son definidas desde el superior al inferior, mientras que los enlaces de las etiqueta se comunican de abajo a arriba. En MPLS, la distribución de etiquetas para la etiqueta de unión se puede lograr a través de un protocolo de distribución de etiquetas LDP.

5.4. CR-LDP (Constraint-based Routing – Label Distribution Protocol) Protocolo de distribución de etiquetas de ruta restringida.³⁰

Este protocolo ofrece características de Ingeniería de tráfico en MPLS, haciendo posible negociar una ruta especial “explícita” para la comunicación, permitiendo establecer LSP’s (*Label Switched Path*) punto a punto con QoS, es llamado

³⁰ DONOSO M. Yesid. Extensión de los métodos Hop-by-Hop, CR-LDP, y RSVP-Te para Multicast IP sobre MPLS. 2001. <http://eia.udg.es/~eusebi/papers/clei.pdf>

protocolo de estado sólido, ya que una vez establecida la comunicación ésta se mantiene “abierta” hasta que se indique lo contrario.

El funcionamiento de este protocolo se da teniendo en cuenta las siguientes particularidades:

- a. Un mensaje LDP_REQUEST permite especificar explícitamente que rutas van a ser utilizadas.
- b. En las rutas explícitas se pueden reservar recursos. Las características del camino pueden ser descritas en *Peak Data Rate* (PDR), *Comitted Data Rate* (CDR), *Peak Burst Size* (PBS), *Committed Burst Size* (CBS), *Frecuency* y *Weight*. PDR y CDR describen las restricciones de ancho de banda de la ruta. PBS y CBS determinan tamaños máximos de las trazas. *Frecuency* especifica que tan frecuentemente se garantiza el CDR. *Weight* determina prioridades relativas en cada LSR para múltiples LSPs cuando hay congestión o poco ancho de banda disponible.
- c. Existe también la opción de indicar que los recursos requeridos pueden ser negociables. Cuando un LSR no puede con los recursos existentes satisfacer un parámetro negociable, realiza la reserva de los recursos disponibles y propaga el LDP_REQUEST con los nuevos valores (evidentemente será un valor menor).

El proceso llevado a cabo por este protocolo está resumido de la siguiente manera:

1. El LER de entrada quiere establecer un nuevo LSP hacia el LER de salida. Los parámetros de tráfico determinan por dónde debe pasar la ruta, así que el LER de entrada reserva los recursos que necesita y envía un mensaje LABEL_REQUEST con la ruta explícita hacia el LER de salida y con los parámetros de tráfico que requiere la sesión.
2. Cada nodo de la ruta que recibe el mensaje reserva los recursos y determina si es la salida para ese LSP, si no lo es, sigue enviando el mensaje LABEL_REQUEST eliminándose de la ruta. Puede reducir la reserva si los parámetros de tráfico están marcados como negociables. (Ver Gráfica 18).
3. Una vez llega al LER de salida, éste realiza cualquier negociación final sobre los recursos y hace la reserva. Asigna una nueva etiqueta al nuevo LSP y la

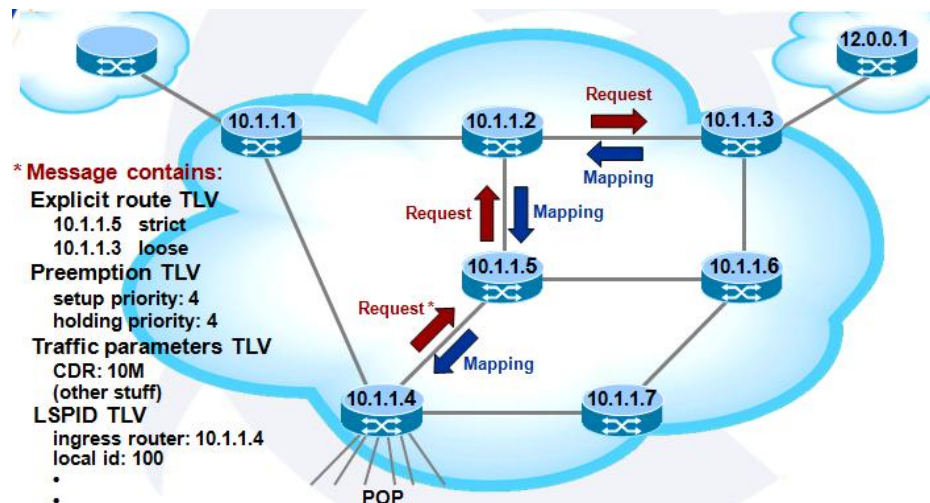
distribuye en un mensaje LABEL_MAPPING que contiene los parámetros de tráfico finales reservados para el LSP. (Ver Gráfica 18).

4. Los LSRs intermedios emparejan los mensajes LABEL_REQUEST y LABEL_MAPPING que han recibido según el identificador de LSP, asignan una etiqueta para el LSP, rellenan la tabla de envío y envían la nueva etiqueta en otro mensaje LABEL_MAPPING.

5. En cuanto llegue al LER de entrada se habrá establecido el LSP. (Ver Gráfica 19).

Gráfica 18

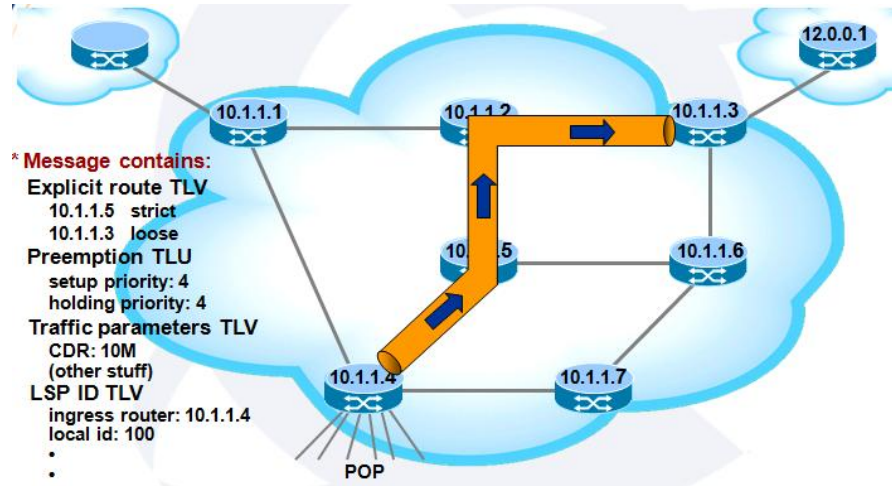
Proceso para determinar ruta y reserva de recursos



Tomada de:

www.itu.int/itudoc/itu-t/workshop/optical/s07-p03r_pp7.ppt

Gráfica 19
Establecimiento de un LSP válido



Tomada de:
www.itu.int/itudoc/itu-t/workshop/optical/s07-p03r_pp7.ppt

5.5. RSVP-TE (Resource reSerVation Protocol – Traffic Engineering): Protocolo de reserva de recursos con ingeniería de tráfico.

Recientemente, el Protocolo de reserva de recursos con la extensión de Ingeniería de Tráfico (RSVP-TE) se ha convertido en el protocolo de distribución de etiquetas para ingeniería de tráfico. Ha sido desarrollado para establecer LSPs explícitos. Existen diferencias entre RSVP-TE y el original RSVP. Por ejemplo RSVP-TE es usado para la señalización entre los routers para establecer los flujos de LSP, a diferencia del original RSVP, el cual se utiliza entre host para crear microflujos. Debido a esto RSVP-TE no tiene el problema de escalabilidad que se presenta en RSVP

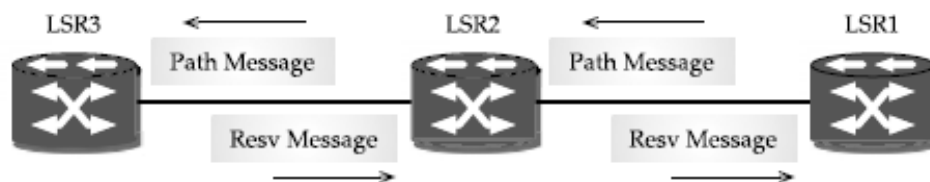
respecto a la gestión de microflujos. RSVP-SE TE utiliza para establecer LSP unicast direccional, mientras que RSVP permite la configuración del flujo de multidifusión.

Grandes flujos de Tráfico, se puede transportar en RSVP-TE por LSP definidos. Un gran flujo de tráfico puede ser dividido en dos LSP establecidos por un LER de entrada y un LER de salida, ó múltiples flujos de tráfico pueden ser combinados para ser transportados por un solo LSP. Por tanto cada LSP sirve como túneles de ingeniería de tráfico. Tenga en cuenta que RSVP-TE no dicta cómo decidir cuándo crear flujos de tráfico diferentes, o cuando dividir flujos de tráfico en múltiples proveedores de red. Más bien, el papel de RSVP-TE es el de un facilitador desde

el punto de vista funcional, mientras que la decisión se deja a los proveedores de operacionales de la red.

Gráfica 20

Distribución de etiquetas usando RSVP-TE



Tomada de: Network Routing: Algorithms, Protocols and Architectures. Pág. 655

En la gráfica 20 se quiere establecer el túnel de ingeniería de tráfico LSR1-LSR2-LSR3. Aquí el mensaje generado en LSR1 se renueva en LSR2 para llegar al destino LSR3. Este ultimo como destino genera el mensaje de RESV en la dirección contraria a LSR2 y que a su vez se remite a LSR1. Así, en RSVP-TE, el mensaje de RESV entonces ayuda a cumplir con la etiqueta de la función de enlace.

El proceso llevado a cabo por este protocolo esta resumido de la siguiente manera:

1. El LER de entrada quiere establecer un nuevo LSP hacia el LER de salida. Los parámetros de tráfico determinan por dónde debe pasar la ruta, así que el LER de entrada envía un mensaje PATH con la ruta explícita hacia el LER de salida y con los parámetros de tráfico que requiere la sesión.
2. Cada nodo de la ruta que recibe el mensaje determina si es la salida para ese LSP, si no lo es, sigue enviando el mensaje PATH eliminándose de la ruta. En cualquier caso cada LSR creará una nueva sesión.

3. Una vez llega al LER de salida, éste determina qué recursos ha de reservar y devuelve un mensaje RESV que distribuirá la etiqueta que ha elegido para ese LSP y contendrá los detalles de la reserva.
4. Los LSRs intermedios emparejan los mensajes PATH y RESV que han recibido según el identificador de LSP, reservan los recursos que indica RESV, asignan una etiqueta para el LSP, rellenan la tabla de envío y envían la nueva etiqueta en otro mensaje RESV.
5. El LER de entrada, cuando lo recibe, enviará un mensaje de confirmación RESVConf para indicar que se ha establecido el LSP.

Después de haberse establecido el LSP se enviarán mensajes periódicos para mantener el camino y las reservas.

PARTE II:
DESCRIPCION DEL PROBLEMA

6. PORQUÉ Y COMO APLICAR INGENIERIA DE TRÁFICO A LAS REDES IP³¹

A mediados de la década de los 90 surgió la necesidad de utilizar Ingeniería de tráfico en las redes debido a las exigencias por información cada vez mayor por parte de los usuarios. En un principio las redes IP no fueron diseñadas para soportar grandes flujos de información de manera simultánea (audio, datos, voz y video). Por tanto el aumento en la demanda de nuevos servicios (Videostream, VozIP, Audiostream, etc) que exigen mayores recursos de red, hizo necesario buscar los mecanismos que garantizaran la transmisión con calidad de toda esa información aprovechando al máximo la infraestructura y los recursos de red existentes. Por tal razón el aumentar el ancho de banda de un enlace de red no resulta ser la mejor alternativa a la hora de establecer en conjunto políticas para la optimización de los recursos disponibles de una red.

En 1996 nacieron los conceptos de conmutación IP. En 1997 la IETF reconoció y empezó a investigar y perfeccionar el concepto de MPLS. En 1999 se estandariza el concepto de ingeniería de tráfico. Así las redes MPLS (Ver Capítulo IV) coexistían a su vez con las redes IP, las cuales no garantizaban calidad en sus servicios debido a que sus técnicas de enrutamiento no habían sido mejoradas, permitiendo que para el año 2000 se determinara los pesos de los enlaces bajo OSPF/IS-IS aplicando así Ingeniería de tráfico a dichas redes, lo que conllevó beneficios a los ISPs (Proveedores de Servicios de Internet) y a los usuarios en general.

Es por ello que el papel de la ingeniería de tráfico es lograr la optimización de una red operativa, de manera que se cumplan los niveles de desempeño y los recursos de red sean bien utilizados. La Ingeniería de tráfico aborda los objetivos de una red, dimensionados en el tiempo según el crecimiento de la misma en usuarios, volumen de tráfico así como el comportamiento general de las redes operativas IP. Normalmente la ingeniería de tráfico no se refiere a la adición de nueva capacidad, la cual se especifica desde el dimensionamiento de la red, sino que hace referencia a optimizar los recursos ya existentes. Además la ingeniería de tráfico no se ocupa de cuestiones tales como el aumento de tráfico surgido de unos pocos segundos a unos pocos minutos que puede resultar en demoras excesivas para un breve período, sino que tiene en cuenta la estadísticas de flujo de información generado por los usuarios de la red en las operaciones que realizan

³¹ MEDHI Deeppankar, RAMASAMY Karthikeyan. Network Routing: Algorithms, Protocols and Architectures. Estados Unidos de América. Morgan Kaufmann Publications. 2007. Pág. 227-244, 676

bajo condiciones normales de funcionamiento. Esto es importante tenerlo en cuenta para comprender el contexto de la ingeniería de tráfico.

6.1. Qué es el tráfico en redes IP?

Para describir el tráfico se debe considerar que existen unos orígenes que generan dicho tráfico IP.

Una red IP provee muchos servicios tales como: Web y E-mail; también existen servicios interactivos tales como Telnet y ssh para servicios de terminales. En las actuales redes IP el tráfico predominante es causado por las aplicaciones que utilizan TCP para la capa de transporte; Se conoce que en un enlace troncal aproximadamente el 90% del tráfico está basado en TCP. El contenido de los mensajes creados por las aplicaciones se divide en pedazos más pequeños de TCP, llamados segmentos TCP, incluyendo la información de encabezado TCP; que luego son transmitidos a través de la red IP, la entidad de datos a nivel de IP es datagramas IP o paquetes. Así, el tráfico en una red IP está compuesto por datagramas IP generados por varias aplicaciones sin preguntarse cuáles son las aplicaciones que los generan.

Al hablar a cerca del volumen de tráfico existente en un enlace de red IP estamos interesados en conocer el número de paquetes IP que fluye en un enlace en una unidad de tiempo determinado. Por tanto el volumen de tráfico puede estar especificado en términos de paquetes IP ofrecidos por segundo. Existe otra medida de volumen de tráfico que se utiliza a menudo como la tasa de transmisión en Megabits por segundo (Mbps) o Gigabits por segundo (Gbps). Determinado por la siguiente ecuación:

Velocidad tráfico de datos = Paquetes por Segundo x Tamaño Promedio del Paquete. (Ecuación 1).

6.2. Tráfico y medidas de desempeño.

En un entorno de red IP, el retardo o demora es un elemento crítico puesto que se debe tener en cuenta para asegurar que un paquete generado por uno de los extremos llegue al otro extremo tan pronto como sea posible. Existe una analogía

entre las redes de transporte por carretera y las redes IP. En las redes de transporte por carretera, la demora depende del volumen de tráfico, así como el número de carriles de la vía y el límite de velocidad, las cuales son impuestas por el sistema. Del mismo modo, el retraso en una red IP puede depender de la cantidad de tráfico, así como la capacidad del sistema; Por tanto se plantea la siguiente relación funcional:

Retardo = F (Velocidad de Volumen de tráfico, Capacidad). (Ecuación 2).

La anterior función solo es factible para un sistema de enlace sencillo. Cuando se considera una red donde el enrutamiento es también un factor importante, entonces surge la siguiente función:

Retardo=F (Velocidad de Volumen de Tráfico, Capacidad, Enrutamiento). (Ecuación 3).

6.3. Caracterización del Tráfico.

El volumen de tráfico puede ser dado a través de un solo número, como el número de paquetes por segundo o la cantidad de Megabits por segundo (Mbps).

¿Cómo se obtiene este parámetro?

Se considera entonces la llegada de paquetes a un enlace de red. Si se considera simplemente una sencilla solicitud a una página web, puede parecer que los paquetes IP están llegando de forma determinística: la página es generada por el servidor web, y se divide en segmentos TCP que se envuelven con un encabezado IP, para luego ser transmitidos unos tras otros. Esto visto desde una sencilla sesión web. Sin embargo, en un enlace o vínculo de red muchas sesiones web son activadas, las cuales son solicitadas al azar por un determinado usuario; esto es similar para el resto del tráfico debido a aplicaciones como correo electrónico y así sucesivamente. Así, desde el punto de vista de un enlace de red, si tenemos en cuenta sólo el nivel de IP, el enlace ve la llegada aleatoria de los paquetes; por tanto, no se puede representar la llegada de estos paquetes con un número fijo. Esta cantidad está más bien determinada por la aleatoriedad de la llegada de tráfico. Por tanto, se puede decir que la medida de dicho tráfico está determinada por la media ó Mbps promedio de llegada de paquetes.

6.4. Acerca de las características del comportamiento aleatorio?

Tradicionalmente, se ha supuesto que la llegada de los paquetes al azar es conocido como un proceso que se denomina el proceso de Poisson; y la tasa promedio de llegadas es la tasa promedio de este proceso. Sin embargo, solo hasta mediados de la década de 1990 hubo una serie de estudios basados en reales

mediciones de tráfico de paquetes que informó que el comportamiento de llegada de paquetes no es Poisson; Por tanto el tráfico es auto-similar en diferentes escalas de tiempo siguiendo pesadas distribuciones de cola. El punto clave para notar aquí es que la auto-similitud contradice el supuesto de Poisson. Además, un proceso auto-similar con una distribución de impactos de cola pesada hace que el comportamiento de los retardos sea peor que un proceso de Poisson. Pero en estudios posteriores se ha observado que de hecho es posible que tanto el comportamiento de Poisson y la auto-similitud pueden co-existir.

Es de resaltar que para objeto de este estudio se tendrá en cuenta el Proceso de Poisson, el cual es el proceso más aleatorio, por lo que no se sabe exactamente el número de llegadas o solicitudes que pueden llegar a un nodo u enlace de la red en un tiempo determinado.

Se debe entender que el ancho de banda se refiere a la velocidad del enlace (Mbps ó Gbps) y el retardo hace referencia al tiempo que tarda la información en viajar por dicho enlace desde el emisor hasta el receptor. Depende de la distancia y del medio.

También es importante desde la perspectiva de la ingeniería de tráfico, que el router tenga suficiente capacidad de memoria, de lo contrario tiende a ser un cuello de botella potencial que conduce a la pérdida de paquetes, reduciendo el rendimiento del protocolo TCP entre host.

6.5. Cuáles son los problemas de una red cuando existe más de un enlace?

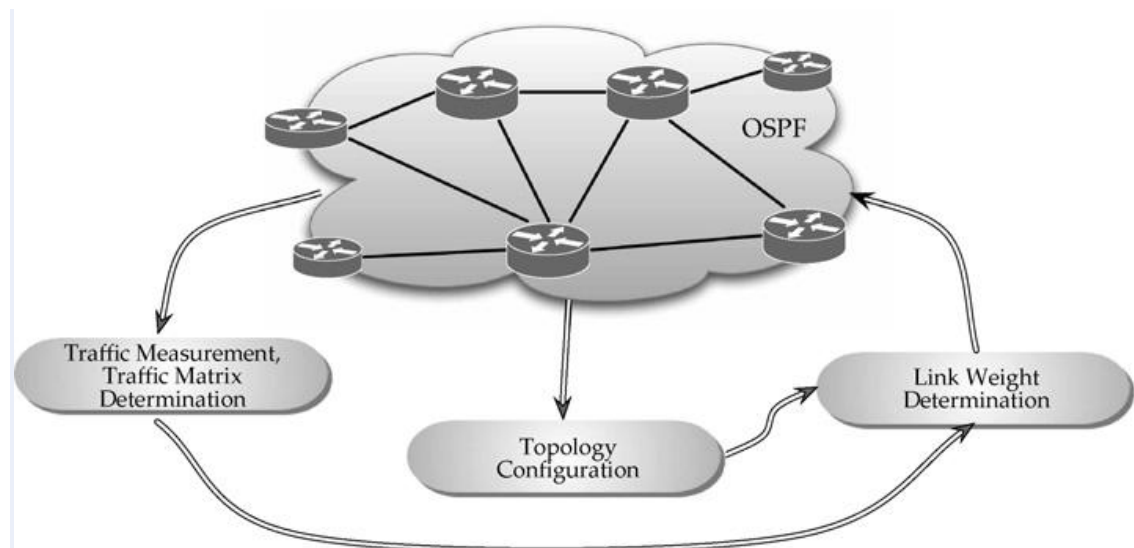
Dado que una red se compone de una serie de routers, existe un volumen de tráfico generado con diferentes demandas y por diferentes pares en la red, además de la capacidad de los enlaces de la red. Por tanto el primer objetivo de la ingeniería de tráfico es optimizar una función objetivo para obtener el sistema

óptimo de pesos de los enlaces, al tiempo que reconoce que la red utiliza el enrutamiento del camino más corto para la transmisión de tráfico.

La descripción anterior requiere un poco más de claridad a la luz de los protocolos OSPF y IS-IS, para saber dónde y cómo encaja la ingeniería de tráfico.

En primer lugar, la ingeniería de tráfico se produce fuera de la red real. Esto puede ilustrarse a través de un marco arquitectónico del sistema de ingeniería de tráfico, como se muestra en la gráfica 21.

Grafica 21
Ingeniería de Tráfico IP – Marco Arquitectónico



Tomada de: Network Routing: Algorithms, Protocols and Architectures. Pag 236

Inicialmente, las mediciones de tráfico se recogen de la red para estimar la matriz de tráfico. Por otra parte, la topología y la configuración también se obtienen de la red. Sobre la base de la topología y la configuración, junto con la matriz de tráfico, se establece un proceso para determinar el peso de los enlaces teniendo en cuenta que OSPF / IS-IS utiliza el enrutamiento del camino más corto. El peso del enlace calculado para cada camino se registra en la red, es decir, cada router recibe una métrica para sus enlaces salientes a través de este proceso externo. Una vez que un enrutador recibe estos parámetros de enlace, a continuación, difunde a los otros routers, a través de inundaciones, los anuncios del estado de

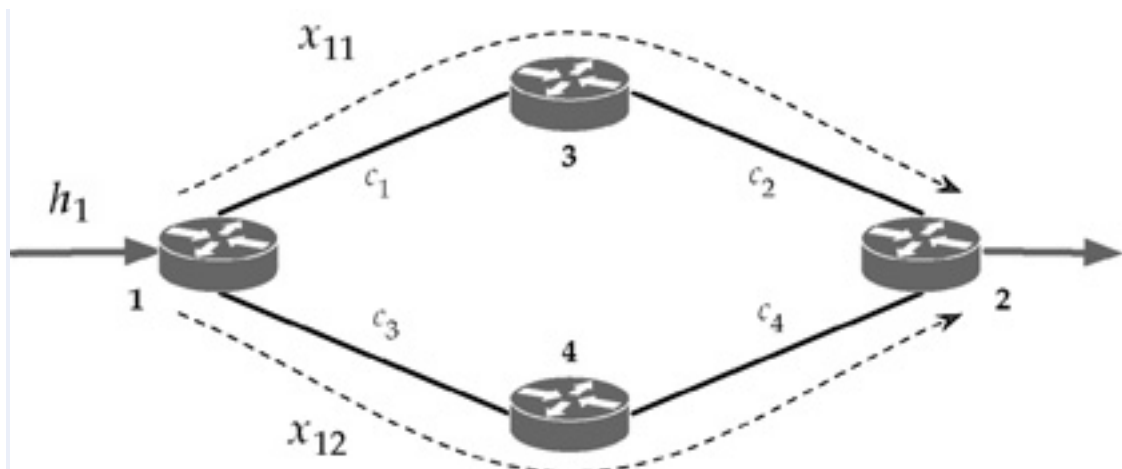
los enlaces (LSA – Link State Advertisements) utilizando los protocolos OSPF/IS-IS.

Esto significaría que si los pesos de los nuevos enlaces no son obtenidos del sistema de ingeniería de tráfico cuando el campo edad de un LSA expira, el router generará un nuevo LSA para continuar el uso del valor métrico del enlace, el cual recibió del último sistema de ingeniería de tráfico. Entonces surge una pregunta: ¿Con qué frecuencia el sistema de ingeniería de tráfico actualiza las ponderaciones del enlace? Esto sin duda correspondería a cada proveedor de la red. En la actualidad la mayoría de proveedores de la red utilizan este tipo de enfoque para actualizar los pesos de los enlaces sea una vez al día o una vez a la semana, en parte para evitar el tráfico de las fluctuaciones a corto plazo, modificando pesos de enlaces con demasiada frecuencia y de otra parte porque la determinación exacta de las medidas de la matriz de tráfico es un proceso bastante complejo y laborioso.

6.6. Optimización del flujo de red

Considere la demanda de tráfico entre los pares 1:2. (Ver gráfica 22)

Gráfica 22
Ejemplo para red de 4 nodos



Tomada de: Network Routing: Algorithms, Protocols and Architectures. Pag 237

El camino desde el nodo 1 hasta el nodo 2 vía nodo 3 se denota como camino 1. El camino desde el nodo 1 hasta el nodo 2 vía nodo 4 como camino 2. Se denotan las variables de flujo como X_{11} y X_{12} , respectivamente. Por tanto para llevar el volumen de tráfico h_1 desde el nodo 1 hasta el nodo 2 se debe cumplir que:

$$X_{11} + X_{12} = h_1. \text{ (Ecuación 4).}$$

También se requiere que los flujos en cada camino no sean negativos

$$X_{11} \geq 0, X_{12} \geq 0. \text{ (Ecuación 5).}$$

Los enlaces son identificados como: 1 para nodos 1-3, 2 para 3-2, 3 para 1-4 y 4 para 4-2. Entonces, podemos enumerar los flujos para satisfacer las limitaciones de la capacidad de la red de la siguiente manera:

$$X_{11} \leq C_1, X_{11} \leq C_2, X_{12} \leq C_3, X_{12} \leq C_4. \text{ (Ecuación 6).}$$

Tenga en cuenta que podemos combinar las limitaciones $X_{11} \leq C_1, X_{11} \leq C_2$ a una única restricción:

$$X_{11} \leq \min \{C_1, C_2\}. \text{ (Ecuación 7).}$$

Esto es similar para las otras dos restricciones; sin embargo se mencionan todas las restricciones para la representación general, a menos que se consideren los valores específicos de la capacidad.

El objetivo es reducir al mínimo la utilización máxima de enlaces, por tanto se puede escribir la optimización del problema como:

$$\begin{aligned} & \text{Minimizar } \{x\} \\ & F = \max \{X_{11}/C_1, X_{11}/C_2, X_{12}/C_3, X_{12}/C_4\}. \end{aligned} \text{ (Ecuación 8).}$$

Sujeto a:

$$X_{11} + X_{12} = h_1 .$$

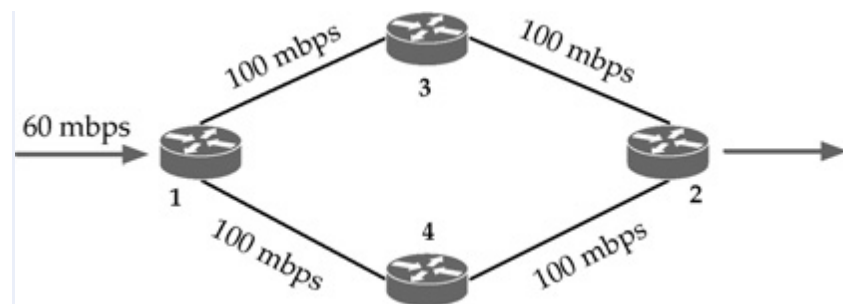
$$X_{11} \leq C_1, X_{11} \leq C_2, X_{12} \leq C_3, X_{12} \leq C_4$$

$$X_{11} \geq 0, X_{12} \geq 0 . \text{ (Ecuación 9).}$$

Ejemplo 1: Todos los enlaces tiene la misma capacidad. (Ver Gráfica 23)

Gráfica 23

Red con 4 nodos y la misma capacidad de los enlaces



Tomada de: Network Routing: Algorithms, Protocols and Architectures. Pag 238

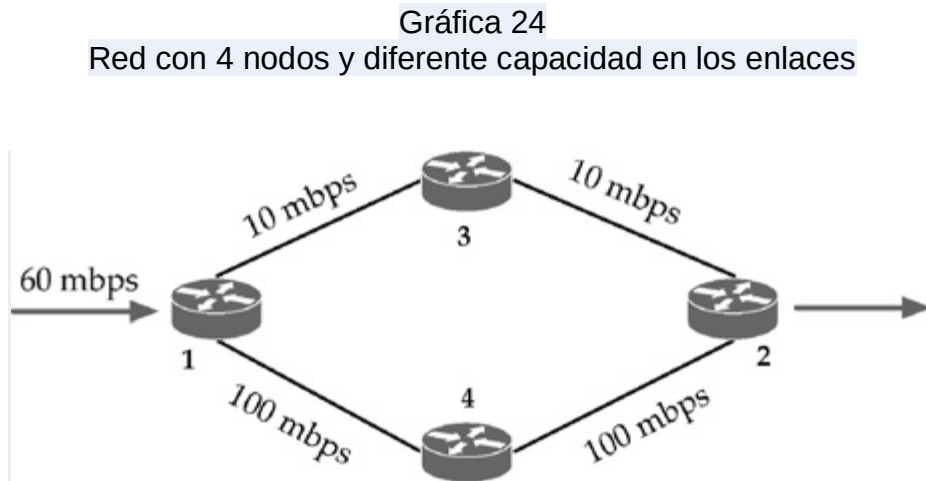
Para el ejemplo todos los enlaces tienen una capacidad de 100Mbps y el volumen de tráfico generado es $h_1=60\text{Mbps}$. El problema se plantea de la siguiente forma:

Minimize [x,r]	$F=r$
Restricciones	$X_{11}+X_{12}=60$
	$X_{11} \leq 100r$
	$X_{12} \leq 100r$
	$X_{11}, X_{12} \geq 0 . \text{ (Ecuación 10).}$

A simple vista se puede observar que la demanda óptima para cada camino sería $X_{11}^* = X_{12}^* = 30$. Al buscar la solución con Programación Lineal se obtiene la utilización máxima del enlace que está dada por $r^* = 30/100 = 0.3$.

Es importante señalar que la utilización máxima de enlaces es una cantidad que oscila entre 0 y 1, ya que el flujo de enlace debe ser inferior a la capacidad del enlace correspondiente para cada problema.

Ejemplo 2: Enlaces con diferente capacidad. (Ver Gráfica 24)



Tomada de: Network Routing: Algorithms, Protocols and Architectures. Pág:239.

Para el ejemplo el volumen de tráfico demandado es $h_1 = 60$, pero la capacidad del enlace para el camino 1 – 3 - 2 es de 10Mbps, mientras que para el camino 1 – 4 - 2 será de 100Mbps. Por tanto resulta mejor enviar más tráfico por el camino 1 – 4 - 2 que por el camino 1 – 3 - 2. Para obtener el equilibrio óptimo se busca la solución óptima, la cual está dada por:

$$X_{11}^*/10 = 60 - X_{11}^*/100$$

$$100X_{11} = 600 - 10X_{11}$$

$$\begin{aligned}
X_{11} + 10X_{12} &= 600/10 \\
11X_{11} &= 60 \\
X_{11} &= 6/11 \\
X_{11} &= 0.545 && \text{(Ecuación 10).} \\
X_{12} &= 60 - X_{11} \\
X_{12} &= 59.45
\end{aligned}$$

Lo anterior implica que la ruta de acceso 1 – 3 – 2 se le asigna un flujo de $X_{11}^* = 0.545$. Mientras que para el camino 1 – 4 – 2 se le asigna un flujo de $X_{12}^* = 60 - X_{11}^* = 59.45$. Por tanto la capacidad máxima del enlace sería: $r^* = 59.45/100 = 0.5945$.

6.7. Enrutamiento por el Primer Camino más corto y Flujo de Red

Una red IP está basada en los protocolos OSPF o IS-IS, donde el camino más corto se calcula en función del peso de enlace (costo o métricas) que se intercambian a través de las inundaciones.

Es importante señalar que este cálculo no tiene en cuenta el volumen de tráfico o la capacidad de la red. Entonces la pregunta que surge es ¿Cómo es el enrutamiento del camino más corto en relación con el modelamiento del flujo de red?

La explicación se hace teniendo en cuenta el ejemplo de la red con 4 nodos, para lo cual los pesos de cada uno de los enlaces se denotarán como w .

Ejemplo 1: La decisión del flujo óptimo con el enrutamiento de la ruta más corta.

Se considera de nuevo el caso en donde todos los enlaces tienen la misma capacidad 100Mbps. Para este caso la optimización de los flujos en la red consistió en dividir el volumen de tráfico en partes iguales para cada camino.

Recordemos que OSPF permite la igualdad de costo-multi-Camino (Equal Cost Multi-Path -ECMP), asignando los mismos pesos para todos los caminos, por lo

tanto, si podemos elegir el peso del enlace de tal manera que se pueda lograr una distribución equitativa del tráfico se podrá lograr el flujo óptimo así como la optimización del flujo en la red. Esto se puede llevar a cabo si damos pesos de enlace = 1 para cada uno de ellos, por tanto el peso de cada ruta será de 2. En otras palabras se trabajará con una métrica basada en saltos.

Qué pasa si los enlaces son de diferente tamaño, es decir, 10 Mbps para los enlaces 1-3 y 3-2 y 100 Mbps para los enlaces 1-4 y 4-2?. Si se conserva todavía la relación de peso a 1 de cada uno, la mitad del volumen total de tráfico de 60Mbps utilizaría la ruta o camino 1 – 3 – 2 debido a ECMP. Sin embargo la capacidad límite sobre este camino es 10Mbps; por tanto un tráfico de 30Mbps ocasionará un gran desbordamiento. Esto significa que debemos utilizar otros pesos en los enlaces para que esto no ocurra. Entonces se requiere que el flujo de tráfico de la ruta 1-3-2 sea menor ya que su capacidad es mucho menor respecto a la capacidad del camino 1 - 4 - 2.

Una manera de lograr esto sería establecer el peso de enlace como la inversa de la relación de la capacidad, es decir a mayor peso del enlace menor capacidad del mismo y viceversa.

Ejemplo:

$$W_1 = W_2 = 1/10$$

$$W_3 = W_4 = 1/100$$

(Ecuación 11).

Lo anterior implica que el costo de la ruta 1-3-2 es 2/10, mientras que el costo de la ruta 1 -4 - 2 es 2/100; Por tanto todo el volumen del tráfico 60Mbps se encaminaría por la ruta 1 - 4 - 2. De hecho la ruta más corta resultó ser la mejor y la máxima utilización del enlace es de $60/100=0.6$.

Dado lo anterior se deduce que los pesos de enlaces actúan como corrientes de flujo. Se observa que X_{11} depende solamente de los pesos de los enlaces W_1 y W_2 , los valores actuales de los otros pesos de enlace W_3 y W_4 son también importantes en la asignación del flujo para X_{11} . Esto es similar para X_{12} . Es decir, los flujos de X_{11} y X_{12} en realidad dependen de todos los pesos en relación W_1 , W_2 , W_3 , W_4 . Si usamos $W = (W_1, W_2, W_3, W_4)$ que indique el vector de pesos de enlace para todos los caminos, entonces podemos escribir la dependencia como

$X_{11}(W)$ y $X_{12}(W)$. Dado que el flujo total debe ser igual al volumen de tráfico demandado, se puede escribir:

$$X_{11}(W) + X_{12}(W) = h_1. \text{ (Ecuación 12).}$$

Luego los requerimientos para cada flujo de enlace en relación a los pesos y a la capacidad de dicho enlace se denotan así:

$$X_{11}(W) \leq C_1, X_{11}(W) \leq C_2, X_{12}(W) \leq C_3, X_{12}(W) \leq C_4$$

(Ecuación 13).

En lo relacionado con el sistema de pesos W hay algunas restricciones sobre qué valores métricos pueden tomar los enlaces; Por ejemplo para OSPF el rango está entre 1 hasta 216 -1, mientras que en IS-IS el rango es de 0 hasta 63.

Mientras que en la demostración de arriba se eligió la métrica del enlace como el inverso de la capacidad del enlace podemos utilizar un factor de normalización para cambiar la métrica del enlace a un rango aceptable.

Por ejemplo, si multiplicamos por 100, entonces la métrica de un enlace con velocidad de 10 Mbps sería de $(100/10=10)$; mientras que la métrica de un vínculo con una velocidad de 100 Mbps sería $(100/100=1)$.

Para simplificar se denota el conjunto de valores permitidos para las métricas de los enlaces como (W) y obtenemos:

Minimize[w,r]: $F=r$

Restricciones:

$$\begin{aligned}
X_{11}(W) + X_{12}(W) &= h_1 \\
X_{11}(W) &\leq C_1 r, X_{11}(W) \leq C_2 r, \\
X_{12}(W) &\leq C_3 r, X_{12}(W) \leq C_4 r \\
X_{11}(W) &\geq 0, X_{12}(W) \geq 0 \\
W_1, W_2, W_3, W_4 &\in W
\end{aligned}
\tag{Ecuación 14}.$$

Hasta el momento, se han utilizado dos reglas simples para la elección de los pesos de enlace. Estas reglas son:

Regla 1: Selección de pesos de enlace que se basa en número de saltos, que se conoce como una métrica basada en saltos.

Regla 2: Selección de pesos de enlace que se basa en la inversa de la capacidad del enlace, que será denominado como métrica basada en el inverso de la capacidad del enlace.

Ejemplo 2: Cambiando el volumen de tráfico

Se considera reducir el volumen del tráfico del par 1:2 de 60Mbps a 5Mbps. (Ver Figura 22). En este caso, se puede utilizar la regla 1 puesto que la capacidad del enlace lo permite (10Mbps) sin enfrentarnos al problema del desbordamiento.

Esto haría que la relación máxima de utilización sea 0,25 (2.5/10) de 2,5 Mbps de volumen de tráfico que se asignarían a la ruta 1-3-2 y el resto de tráfico (2.5Mbps) es asignado a la ruta 1-4-2 debido a ECMP obteniéndose una relación máxima de utilización de 0.025(2.5/100). Si se utiliza la regla 2, entonces todo el tráfico se distribuirá por la ruta 1 – 4 - 2, por tanto, la relación de máxima de utilización del enlace es: 0,05 (5/100). Con cualquier regla, podemos ver que cuando el volumen de tráfico es bajo en comparación con la capacidad de la red, la utilización sigue siendo baja.

Ahora se considerará la topología con todos los enlaces a 100Mbps. (Ver Figura 16)

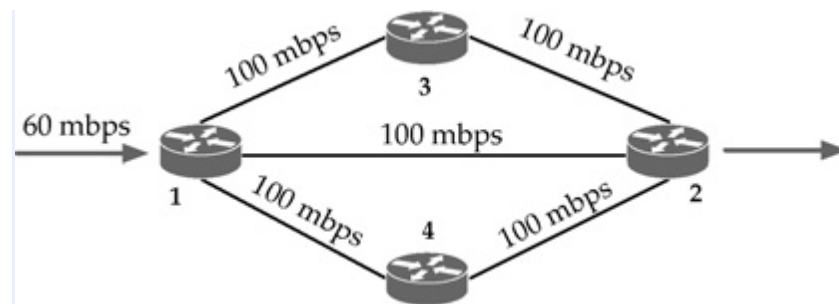
Para este caso, con 5 Mbps de tráfico de red, se llega a la misma máxima utilización del enlace 0,025 ($2.5/100$) para cada uno de los caminos con la regla 1 para el peso de enlaces, en total sería ($0.025*2=0.05$). Por tanto si todos los enlaces de una red tienen la misma capacidad y/o velocidad entonces el peso del enlace basado en el inverso de la capacidad de conexión (Regla 2) es el mismo que el de la métrica basada en los pesos de los enlaces ($5/100=0.05$). Así la regla 1 se puede considerar como un caso especial de la regla 2.

A continuación se considera un ejemplo donde, debido a un aumento previsto del tráfico de la red un nuevo enlace, es añadido. Esto en realidad hay que tenerlo en cuenta en el problema de dimensionamiento de la red.

Ejemplo 3: Cambio de topología en previsión al aumento en el volumen de tráfico

Gráfica 25

Red de 4 nodos con 5 enlaces



Tomada de: Network Routing: Algorithms, Protocols and Architectures. Pág:242.

En esta gráfica se puede apreciar un nuevo enlace directo entre los nodos 1 y 2. Este es el enlace número 5, cuyo peso es W_5 , los demás enlaces siguen siendo los mismos. Este enlace ha sido añadido en previsión de un aumento en el volumen de tráfico. Ahora tenemos otro camino posible entre los nodos 1 y 2. Este camino se le asocia la variable de flujo X_{13} .

Entonces tenemos:

Minimizar $[w,r]$

$F=r$

Restricciones

$$X_{11}(W) + X_{12}(W) + X_{13}(W) = 60$$

$$X_{11}(W) \leq 100r$$

$$X_{12}(W) \leq 100r$$

$$X_{13}(W) \leq 100r \quad (\text{Ecuación 15}).$$

$$X_{11}(W), X_{12}(W), X_{13}(W) \geq 0$$

$$W_1, W_2, W_3, W_4, W_5 \in W$$

Podemos ver fácilmente que, independientemente de si usamos la regla 1 o la Regla 2 todo el tráfico se destinará al enlace número 5 por ser la ruta más corta al ser un enlace directo. En este caso la máxima utilización del enlace es $r = 60 / 100 = 0.6$.

De lo anterior se puede ver que la regla 2 no siempre funciona bien ya que en este caso mediante la adición del nuevo enlace re-direccionamos todo el tráfico hacia él, en lugar de la misma asignación de flujo entre los tres caminos. Esto significa que se ha aumentado la utilización máxima del enlace, aumentando así el retardo promedio de la red para un volumen de tráfico de 60Mbps mediante la adición de un nuevo enlace, en comparación cuando la relación directa no existe.

Contrariando la lógica o intuición, este nuevo enlace que surge ante la necesidad de dimensionar la red debido al aumento en el volumen de tráfico y el cual posee la misma capacidad de los otros dos, no resulta ser el más efectivo ya que todo el flujo de información se direccionaría por este al ser la ruta más corta y con menor peso total, caso mencionado en la paradoja de Braess³², esto puede suceder en redes IP, si los pesos de los enlaces no se eligen correctamente.

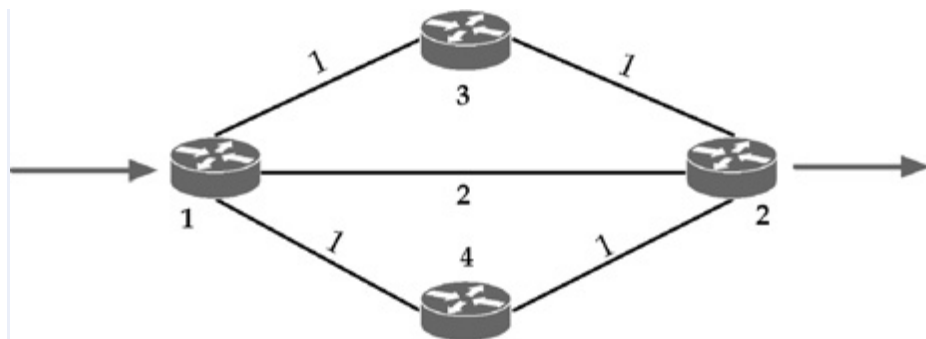
En la red de la gráfica 25 surge la pregunta se puede escoger un mejor conjunto de pesos de enlaces que reducirá la utilización máxima del enlace con respecto a cuándo se utiliza la regla 2? Es posible encontrar una mejor solución?

³² LOZANO, Angelica. TORRES Vicente y ANTÚN Juan Pablo. Tráfico Vehicular. 2003. [Http://www.ejournal.unam.mx/cns/no70/CNS07004.pdf](http://www.ejournal.unam.mx/cns/no70/CNS07004.pdf).

Esto es posible si elegimos los pesos enlace de la siguiente manera (Ver Gráfica 26).

Gráfica 26

Red de cuatro nodos con pesos de enlace



Tomada de: Network Routing: Algorithms, Protocols and Architectures. Pág: 243.

$$W_1 = W_2 = W_3 = W_4 = 1$$

$$W_5 = 2$$

(Ecuación 16).

De esta manera, el coste de la ruta para todos los caminos son 2; Así, debido a ECMP, el volumen de tráfico será igualmente dividido entre los tres caminos (60Mbps/3=20Mbps para cada camino) reduciendo así la utilización máxima de enlaces, r , (20Mbps/100Mbps=0.2) en total (0.2*3=0.6). De hecho este conjunto de pesos de enlace son óptimos para el problema planteado. Aunque esta relación de pesos no son los únicos a elegir. Por ejemplo si se escoge

$$W_1 = W_2 = W_3 = W_4 = 2$$

$$W_5 = 4 \text{ (Ecuación 17).}$$

El resultado será el mismo valor para la utilización máxima del enlace, (20Mbps/100Mbps=0.2) esto para cada camino, luego (0.2*3 caminos= 0.6).

PARTE III:

PROPUESTAS EXISTENTES

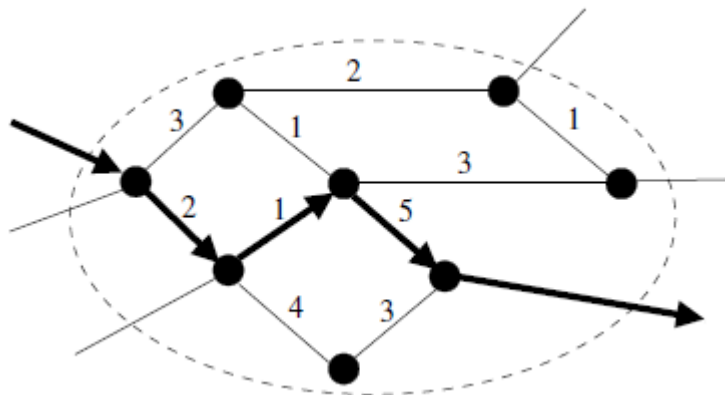
7. INGENIERIA DE TRÁFICO CON PROTOCOLOS DE ENRUTAMIENTO IP TRADICIONALES³³

Otros estudios confirman que la mayoría de las grandes redes IP corren protocolos IGP's (Interior Gateway Protocols) como el OSPF ó el IS-IS que seleccionan caminos basados en los pesos de enlace estático. Tales pesos son típicamente configurados por operadores de la red.

Los routers usan estos protocolos para el intercambio de pesos en los enlaces y construyen una completa vista de la topología dentro de un AS (Autonomous System). Cada router calcula el camino más corto (la longitud del camino es la suma de pesos de los enlaces) y crea una tabla que controla la transmisión de cada paquete IP hasta el siguiente salto en su ruta. Ver Gráfica 27.

Gráfica 27

Enrutamiento del camino más corto dentro de un Sistema Autónomo basado sobre OSPF/IS-IS con pesos de enlace: cada enlace tiene un peso entero.



Tomada de Artículo: Traffic Engineering With Traditional IP Routing Protocols. De: ScienceDirect 2004. Pág.2.

³³ FORZT, Bernard, REXFORD Jennifer and THORUP Mikkel. Artículo: Traffic Engineering With Traditional IP Routing Protocols. Universite Catholique de Louvain, Belgium; AT&T Labs-Research. Florham Park. De: ScienceDirect 2004.

En sus formas básicas el OSPF / IS-IS no adaptan los pesos de los enlaces en respuesta a los cambios en el tráfico ó en las fallas en los elementos de la red. Los recientes estándares han propuesto extensiones de ingeniería de tráfico para OSPF/IS-IS donde se incorpora carga de tráfico en el estado del enlace y decisiones de selección de ruta. Sin embargo estas extensiones requieren modificaciones de los routers para recoger y difundir las estadísticas de tráfico y establecer caminos basados en las métricas de carga. Pero esto no es necesario si se cuenta con una buena selección de pesos de enlace estático que son resistentes a las fluctuaciones de tráfico y a las fallas de los enlaces permitiendo el uso del tradicional OSPF/IS-IS.

La gráfica 28 muestra cómo obtener el control de la distribución del tráfico en una red seleccionado pesos IGP. Los tres diagramas conciernen a la misma red donde los enlaces (u,t) (v,w) y (w,t) tienen la misma capacidad 2 y cada uno de los nodos q, r, s y w tienen una unidad de tráfico para enviar al nodo t y su capacidad es de 1. El objetivo es minimizar la máxima carga del enlace.

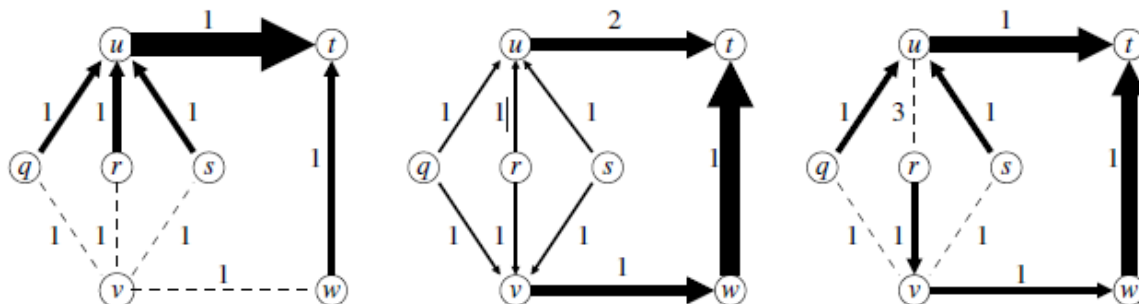
El primer diagrama muestra los resultados de tener el mismo peso de 1 en cada enlace. Todo el tráfico se direcciona desde los nodos q, r y s por el nodo u, forzando a 3 unidades de carga el enlace (u, t).

Una aproximación para reducir la carga es incrementar a 2 el peso IGP en el enlace sobrecargado (u, t). Si se observa el segundo diagrama este cambio en el peso del enlace (u,t) resulta en dos caminos más cortos para los nodos q, r y s e incluso un fraccionamiento del tráfico por los caminos vía u y v de (1,5 c/u). Sin embargo esta solución da lugar a 2,5 unidades de carga transportada por el enlace (w, t).

Una optimización global de los pesos producirá un ajuste en los pesos como en el tercer diagrama, donde no habrá enlace transportando más de 2 unidades de carga o de tráfico. De tal forma que los enlaces al nodo t quedan equilibrados con la misma carga, para un total de carga transportada de 4. Por tanto este resulta ser la mejor solución para el enrutamiento general de la red.

Gráfica 28

Enrutar las mismas demandas con diferentes configuraciones de peso; cada enlace tiene un peso entero, las flechas ilustran el flujo del tráfico, es grosor de la flecha indica el volumen de tráfico atravesando el enlace y la línea punteada representa un enlace que no transporta tráfico.



Tomada de Artículo: Traffic Engineering With Traditional IP Routing Protocols. De: ScienceDirect 2004. Pág.3.

7.1. Ventajas de la utilización tradicional del OSPF / IS-IS

Es claro que el marco existente de pesos de enlace estático resulta ser práctico para trabajar sin tener que modificar los protocolos de enrutamiento o de los router en sí. Es decir que los procesos de llegar a buenos valores para los pesos de enlaces pueden depender de las mediciones de tráfico y de la topología de la red. Dichos pesos también pueden depender de amplia variedad de diferentes métricas tales como: mayor rendimiento, menor retardo, etc.

Los pesos son configurados por una entidad externa, (administrador del sistema de red) el cual determina los objetivos de ingeniería de tráfico que se debe cumplir.

Usar pesos de enlace para expresar la configuración del enrutamiento tiene las siguientes ventajas:

- *Compatibilidad con el tradicional camino más corto IGPs:* trabajar con la existencia de los protocolos de enrutamiento OSPF / IS-IS, evita la necesidad de actualizar equipos ó introducir configuración adicional compleja.

- *Representación Concisa:* los pesos de enlace son una forma precisa del estado de la configuración. Cada router está configurado con el peso para cada uno de sus enlaces salientes. El router no necesita conocer información adicional relacionada con otros routers.
- *Pesos por defecto y backups de rutas:* los pesos de enlace pueden tener una razonable configuración por defecto basada en la capacidad del enlace. Si la topología cambia debido a la falla de un enlace, el router puede automáticamente calcular nuevas rutas basadas en la topología actual y en los pesos de enlace. Estos routers pueden transportar tráfico hasta que la nueva configuración sea seleccionada, si es necesario. Algunos administradores de red ya ajustan sus pesos IGP en respuesta a la congestión de la red.

7.2. Selección de camino basado en Configuración IGP

La ingeniería de tráfico requiere un camino para predecir el flujo de tráfico a través de la red basado en la configuración de enrutamiento. Cuando todos los enlaces pertenecen a una única área OSPF / IS-IS, la ruta seleccionada simplemente implica calcular el camino ó caminos más cortos entre cada par de routers usando el algoritmo Dijkstra.

Redes mas grandes están típicamente divididas en múltiples áreas OSPF / IS-IS. Para routers en diferentes áreas, el camino seleccionado depende de la información transmitida en los límites del área. En algunos casos la red puede tener múltiples caminos cortos entre el mismo par de routers. En la práctica la mayoría de los routers aprovechan los múltiples caminos para balancear la carga. Un router divide el tráfico más o menos uniformemente sobre cada uno de los enlaces salientes a lo largo de un camino más corto hasta el destino.

Tradicionalmente los pesos IGP's de enlace estático no tiene la flexibilidad para dividir el tráfico entre los caminos más cortos en condiciones no previstas. Entonces el modelo de enrutamiento debe calcular un conjunto de rutas para cada par de routers. Estos caminos pueden ser representados en términos de la fracción del tráfico (entre un par de routers) que atraviesa cada uno de los enlaces. Por tanto el modelo de enrutamiento puede ser combinado con la demanda de tráfico para estimar el volumen de tráfico en cada enlace, basado en la topología y la configuración IP.

7.3. Control: Reconfiguración de pesos IGP

Cambiar pesos IGP requiere la aplicación de comandos adecuados para las rutas afectadas. Esto puede implicar usar telnet ó ssh sobre la interfaz de cada router. Estos comandos específicos dependen del sistema operativo que corra en el router. Además pueden ser aplicados manualmente por un operador de red o automáticamente por un script.

Luego de un cambio de peso, el router actualiza su estado de enlace en la base de datos y los nuevos pesos para el resto de la red. Tras la recepción cada router actualiza su base de datos de estados de enlace, calcula el nuevo camino más corto y actualiza ciertas entradas en su tabla de reenvío.

La convergencia en la red tras un cambio de peso es más rápida que la convergencia después de un fallo ya que el router no tiene que incurrir en un retraso para detectar que hubo un error. Aún cambiando los pesos del enlace la red requiere de un periodo de transición donde los caminos o rutas de transmisión están cambiando para parte del tráfico. Como tal no se deben hacer cambios frecuentes para pesos de enlace. Estos cambios deben obedecer a circunstancias especiales como nueva capacidad del enlace o un serio cambio en la demanda del tráfico.

7.4. Propiedades de rendimiento

Se puede aplicar ingeniería de tráfico mediante el uso de los protocolos tradicionales de enrutamiento OSPF / IS-IS obteniendo así una buena configuración de pesos en los enlaces para así satisfacer un rendimiento en particular de la red.

Cuando los enlaces tienen diferentes capacidades se recomienda utilizar la relación entre carga y capacidad del enlace.

La capacidad de un enlace se considera como su máxima carga deseable. Por tanto el objetivo es mantener la máxima utilización por debajo del 100%; en caso de la protección contra ráfagas sería mantener su máxima utilización en un 60%.

7.5. Topologías fijas y demandas de tráfico

Cuando se juzga la calidad de una configuración de pesos, comparamos eso con el OPT (Optimal Routing), el cual puede direccionar el tráfico a lo largo de cualquier camino en cualquier proporción. Modelos OPT's son un esquema de enrutamiento idealizado que puede establecer uno o más caminos explícitos entre cada par de nodos y distribuir cantidades de tráfico de manera arbitraria sobre cada camino.

Realizar una solución OPT en la práctica requiere de un protocolo de enrutamiento más flexible, tal como MPLS, que soporta rutas explícitas. Las propuestas actuales como la de Cisco consiste en configurar pesos de un enlace inversamente proporcional a su capacidad (InvCapOSPF). Otra es la UnitOSPF la cual establece cada peso de los enlaces en 1.

7.6. Comparación del Rendimiento con Máxima utilización del enlace

En la Gráfica 18 los nodos q, r, s y w envían una unidad de tráfico a t. La capacidad de los enlaces incidentes locales para s, r y q es de 1, mientras los enlaces restantes (u, t) (v, w) y (w, t) tienen capacidad de 2. UnitOSPF da la configuración de peso del primer diagrama. La máxima utilización es para el enlace (u,v). Con una carga de 3 y una capacidad de 2, su utilización es de 150%. Al aplicar InvCapOSPF se reducirán los pesos de los enlaces (u,t), (v,w) y (w,t) a la mitad. Esto no afecta el enrutamiento del tráfico desde q, r, y s, ya que sigue siendo mas corto ir a través de u en lugar de v. Así la máxima utilización sigue siendo 150%. Sin embargo la utilización de pesos en el último diagrama logra una utilización máxima del 100%. El 100% de utilización sucede sobre el enlace (u,t) y (w,t) con carga y capacidad 2 y sobre el enlace (q,u) (s,u) y (r,v) con carga y capacidad 1. En este caso OPT no puede realizar mejoras ya que el tráfico total para t es 4 y dicho tráfico es el que recibe por sus dos entradas.

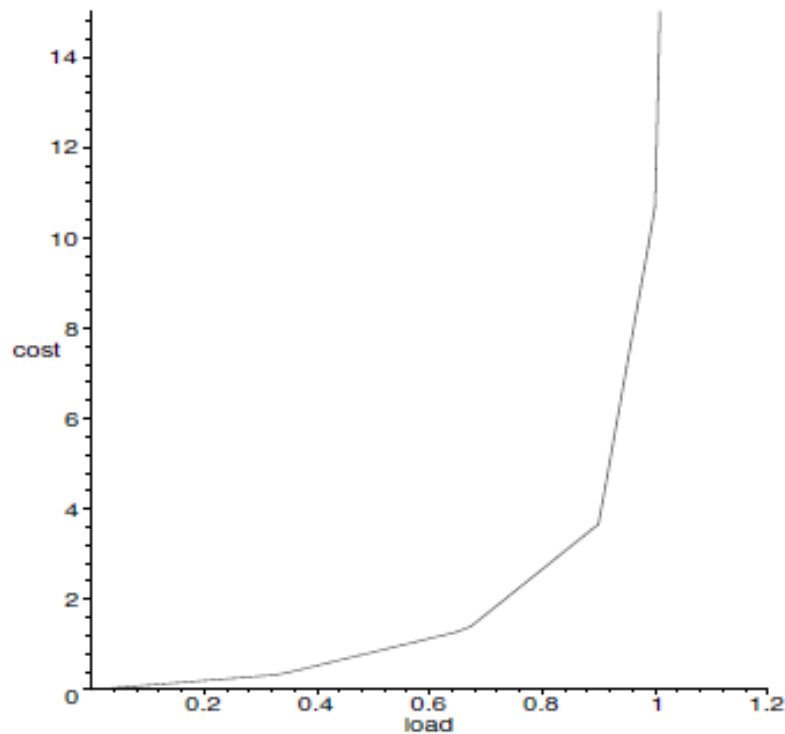
Estudios han demostrado que no siempre los protocolos de enrutamiento flexible son necesarios para un buen balanceo de cargas. Pues la configuración de pesos AdvancedOSPF puede dar buenos resultados. Es decir el punto de partida es el InvCapOSPF ó UnitOSPF para luego llegar al AdvancedOSPF. Esto con el fin de ser capaz de soportar 50% de mayores exigencias, una muy atractiva alternativa para comprar y desplegar enlaces adicionales. Además el AdvancedOSPF configura pesos dejando muy poco espacio para mejoras por parte de cualquier otro protocolo de enrutamiento más flexible.

7.7. Comparación de rendimiento con una red de gran capacidad

Al observar la gráfica 28 si la capacidad de los enlaces fueran 5 veces mas grandes, entonces no valdría la pena desviar el tráfico de r por v. Por tal motivo se considera una red con mayor capacidad. El costo de usar un enlace se incrementa con la utilización o carga, la cual puede ir creciendo hasta exceder el 100%. Una función del costo del enlace se muestra en la gráfica 29. El umbral máximo de carga es 1.

Gráfica 29

Costo del enlace en función de la carga y con capacidad 1



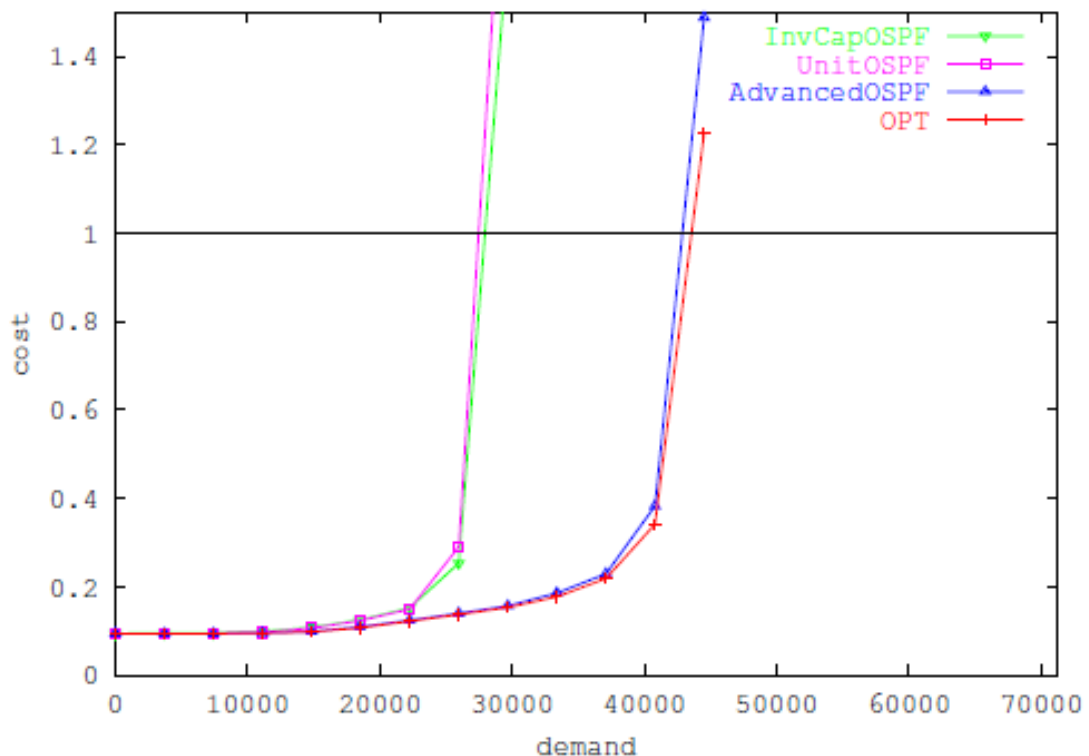
Tomada de Artículo: Traffic Engineering With Traditional IP Routing Protocols. De: ScienceDirect 2004. Pág.10

El AdvancedOSPF puede manejar una demanda mayor hasta del 50% comparada con el InvCapOSPF y UnitOSPF antes de pasar el umbral. (Ver Gráfica 30).

Dicha gráfica muestra el resultado para una propuesta de AT&T backbone – columna vertebral o principal de tráfico de internet basado en mediciones de tráfico y en las tendencias de crecimiento.

Gráfica 30

Costo de Red con mayor capacidad Vrs Demanda para una propuesta de Backbone AT&T



Tomada de Artículo: Traffic Engineering With Traditional IP Routing Protocols. De: ScienceDirect 2004. Pág.11

Como se observa OPT es la solución óptima de enrutamiento general con respecto al costo de la red con mayor capacidad, pero esta puede solamente manejar 2% más demanda que AdvancedOSPF. La curva para AdvancedOSPF se

acerca mucho a OPT cuando la utilización máxima pasa el 100%. Esto se debe a que AdvancedOSPF intenta ajustar el tráfico para utilizaciones por encima del 100% de demanda.

Por tanto la máxima utilización se mantiene por debajo del 100%, AdvancedOSPF puede manejar 50% más de tráfico comparado con InvCapOSPF y UnitOSPF, mientras que OPT puede solamente manejar 3% más que AdvancedOSPF.

7.8. Reduciendo el número de los valores de peso

Para simplificar el proceso, es mejor utilizar un pequeño número de diferentes valores de peso. Además teniendo un conjunto de valores pequeños se reduce significativamente la sobrecarga del algoritmo que explora posibles cambios en los pesos. En los experimentos se recomienda entre 1-20 lo cual resulta ser suficiente para ser competitivo con OPT.

7.9. Cambiando la demanda de tráfico

Optimizar los pesos para una topología simple y para la matriz de tráfico no es suficiente. En la práctica el volumen del tráfico fluctúa con el tiempo y las fallas inesperadas pueden resultar en cambios para la topología de red. Además adquirir un estimado exacto de la matriz es difícil. Es importante que la configuración de los pesos IGP sea robusta para resistir a los cambios en el tráfico y la topología. Para poner a prueba la solidez, evaluamos la configuración optimizada del peso del enlace con diferentes matrices de tráfico.

Una configuración optimizada de los pesos del enlace hace un buen uso de la capacidad de los enlaces de la red y esto no es sensible a pequeños o moderados cambios en el tráfico.

Además las fluctuaciones estadísticas de tráfico cambian algún día. Un ejemplo sería considerar dos matrices de tráfico (diurno y nocturno) para lo cual una solución con pesos IGP optimizados para el tráfico diurno no necesariamente tiene buen desempeño para tráfico nocturno. Los experimentos indican que una configuración de peso simple se puede implementar muy bien para ambas matrices de tráfico, casi también como establecer pesos optimizados para cada matriz. Pero usando la configuración de pesos simples durante todo el día tiene dos ventajas: la primera, los operadores no necesitan interrumpir la red con

cambios en los pesos para cualquier enlace; segundo la configuración del peso común funciona bien para todas las combinaciones de las dos matrices, es decir es eficaz en la acogida de la transición gradual entre el tráfico de día y de noche.

7.10. Pocos cambios para los pesos de enlace

Últimamente, cambiar los pesos de enlace es necesario en respuesta a los grandes cambios en el tráfico y determinadas fallas en los enlaces. Limitar el número de cambios de peso es importante para limitar la interrupción en la red. Afortunadamente, cambiar el peso de un enlace simple es muy eficaz.

Como medida proactiva, los cambios en los pesos deben ser calculados de antemano ante cualquier fallo de enlace y almacenado en el sistema de administración de la red.

Encontrar una buena ruta con un pequeño número de cambios en los pesos sugiere que operadores puedan seguir un enfoque híbrido en la asignación de pesos de enlace. Configurar el peso inversamente proporcional a la capacidad del enlace (InvCapOSPF) tiene la ventaja de ser intuitivo y simple. Experimentos indican que al aplicar InvCapOSPF en 10 de los enlaces de una red fue suficiente para que dicha solución se aproximara a OPT.

Además, algunos protocolos soportan una ruta “pinning” el cual permite el movimiento de parte del tráfico de un camino a otro sin perturbar los caminos restantes. Otros establecen backup de caminos permitiendo mayor enrutamiento en el evento de una falla en la red. Estas dos mejoras tienen el potencial para reducir la interrupción temporal del tráfico existente en caso de cambios de enrutamiento y fallas en la red. A pesar de las ventajas de los más avanzados protocolos de enrutamiento, la simple aproximación de ajustes de pesos de enlace estático puede ser una solución viable en las redes IP.

8. HACIA LA BAJA COMPLEJIDAD DE LA INGENIERÍA DE TRÁFICO EN INTERNET: LA ADAPTACIÓN DEL ALGORITMO MULTI CAMINO³⁴

Recientes estudios describen un nuevo algoritmo para ingeniería de tráfico dinámico en Internet llamado Enrutamiento Adaptativo Multicamino (AMP). El objetivo principal de AMP es la de distribuir la carga dentro de un dominio de red de manera continua mediante la descarga de enlaces congestionados en tiempo real. La estructura de señalización de este algoritmo recursivo representa una nueva contribución, ya que restringe el intercambio de información a los nodos vecinos manteniendo al mismo tiempo la propagación de la congestión de la información en la red.

8.1. Ingeniería de tráfico para enrutamiento con estado de enlace

El enrutamiento del tráfico en las redes IP que utilizan la ruta más corta de transmisión basado en los pesos de enlace representa hoy una de las técnicas concretas implementadas por importantes proveedores de servicios de Internet (ISP) - Service Provider Networks.

Debido a la madurez de los protocolos de estados de enlace, estas redes muestran una robustez muy alta ya que el tráfico es rápido y de fiable distribución en las rutas alternativas ante la presencia de fallas en los equipos lo cual resulta ser la principal preocupación de cualquier proveedor de red. Aparte de estas ventajas, los algoritmos de enrutamiento de estado de enlace tienen el problema de la configuración adecuada de la red, en particular la asignación de ponderaciones de enlace la cual influye sobre los routers en grandes redes. Desafortunadamente, los fabricantes de equipos de enrutamiento IP, como Cisco Systems y otros se han limitado a implementar en la práctica varias mejoras heurísticas para establecer el peso de los enlaces. Una de ellas es el establecimiento de pesos de enlace inversamente proporcional a la capacidad del enlace. Estos enfoques no adaptan la ruta o camino a las condiciones actuales de tráfico en las redes. En muchos casos, esto puede llevar a la congestión en las redes que tienen el potencial de alojar un mayor volumen de tráfico. Por lo tanto, el estado actual de la técnica en el enrutamiento de estado del enlace deja mucho

³⁴ GOJMERAC Ivan, REICHL Peter and JASEN Lassen. Artículo: Towards Low-complexity internet traffic engineering the adaptive Multi-Path Algorithm. De: ScienceDirect 2008.

margen de mejora utilizando métodos de ingeniería de tráfico, cuyo objetivo es optimizar la utilización de sus recursos.

8.2. Avances recientes en enrutamiento de múltiples caminos

En los últimos años, el creciente interés en optimizar el enrutamiento por múltiples rutas ha llevado a varias contribuciones importantes que cubren una amplia gama de investigaciones en diferentes direcciones. Desde una perspectiva teórica, se trata de describir un lenguaje matemático común para el diseño de redes a fin de proporcionar una base sólida para un trabajo más práctico.

Partiendo de un escenario general de enrutamiento multi-camino con rutas predefinidas, un nuevo algoritmo de optimización es propuesto y evaluado. El algoritmo Multi-Camino Optimizado (OMP) es de lejos el más importante aporte al permitir que los protocolos de enrutamiento actuales tengan capacidad para aplicar ingeniería de tráfico sin dejar de ser totalmente compatible con la arquitectura de enrutamiento existente. Con el fin de lograr este objetivo, OMP se basa en las propiedades favorables de enrutamiento de estado del enlace, como el protocolo Open Shortest Path First (OSPF) y el criterio del mejor camino para maximizar el número de múltiples caminos libres. En contraste con los enfoques de equilibrio de carga que actúa dependiendo de la configuración del peso del enlace, OMP ha elegido un enfoque novedoso que permite granularidad muy alta de desplazamiento de carga, allanando así el camino hacia un funcionamiento estable en entornos de red más reales. Por tanto la información global de la carga de los enlaces en las redes es usada para ordenar el tráfico en la ruta de tal manera que se alivia la congestión en los puntos de acceso por medio de la redistribución del tráfico. Considerando que la información de carga de los enlaces es distribuida a través de la red a través del mecanismo de inundación SLA (Link State Advertisement), las acciones de equilibrio de carga se realizan estrictamente a nivel local en cada nodo sin recurrir a ningún tipo de coordinación global.

Basado en la información recibida acerca de la carga del enlace, el algoritmo ajusta dicha carga (esto se conoce como el núcleo de OMP) identifica la más alta, cargando el enlace posterior de la red, y ajusta las cuotas de la carga relativa de tráfico por el siguiente salto y para cada destino de tal forma que la carga es distribuida en el menor número de caminos intermedios. OMP define un mecanismo de balanceo de carga muy sofisticado para controlar la dinámica de desplazamiento de dicha carga, permitiendo alcanzar la distribución del tráfico con el mínimo desbordamiento en el enlace.

El desempeño del algoritmo Multi-camino optimizado es evaluado y comparado con los resultados de dos estrategias de enrutamiento estándar, tales como: el *enrutamiento del camino más corto (SPR)* y el denominado *Igualdad de costo en enrutamiento multi-camino (ECMP)* a través de diferentes investigaciones y simulaciones realizadas por National Science Foundation, NSF.

El rendimiento de OSPF-OMP viene siendo investigado en varios escenarios y se viene comparando con el enrutamiento del camino más corto y ECMP. Resumiendo, los resultados OMP ha demostrado actuar de mejor forma que las estrategias de enrutamiento de la competencia, mostrando mejoras significativas en términos de tiempos de entrega de paquetes y en las tasas de pérdida de paquetes en relación a la carga de tráfico ligero y medio, mientras que las otras dos estrategias de enrutamiento muestran tasas de pérdida notable de paquetes. Por otra parte OMP también ha mostrado una favorable distribución del tráfico, permitiendo así una mejor adaptación ante los cambios repentinos en la demanda del tráfico comparado con las técnicas de enrutamiento más corto (SPR) y ECMP.

Sin embargo, es importante mencionar que la complejidad del algoritmo podría representar una preocupación importante en el contexto del despliegue operativo en las redes IP. Los factores más importantes que contribuyen a la complejidad son los requisitos de grandes cantidades de memoria para el almacenamiento y la gestión de las estructuras del siguiente salto en cada uno de los routers, puesto que contienen múltiples caminos hacia todos los destinos en la red. Aún más importante, el proceso de regulación de la carga es muy complejo, ya que se realizan cálculos por separado para cada estructura multi-camino. Además como factor agravante, la señalización de información de carga usando el algoritmo de inundación de estado de enlace es un proceso indeterminado, lo que dificulta notablemente los esfuerzos por determinar los costos generales en el peor de los casos para el algoritmo.

8.3. Algoritmo de enrutamiento adaptativo multicamino (AMP)

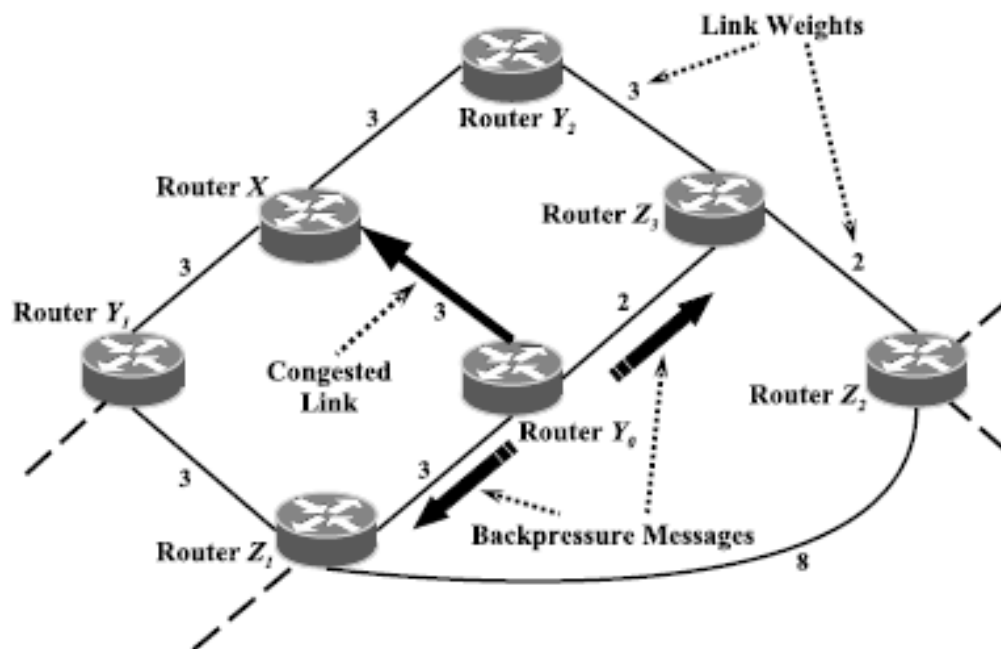
El principal objetivo en el desarrollo de AMP consiste en diseñar una heurística o mecanismo de baja complejidad que permita el balanceo de carga continua en dominios de enrutamiento de forma individual. El algoritmo se ejecuta de manera totalmente autónoma en el plano de enrutamiento de los nodos sin necesidad de ninguna intervención humana.

A fin de mantener ventajas OMP en términos del nivel de automatización y el uso de múltiples estructuras de ruta, y al mismo tiempo tratar de aliviar los

inconvenientes mencionados, principalmente en lo que respecta a las inundaciones mensaje como un medio para la señalización global, AMP introduce una nueva y fundamental diferencia en la arquitectura de señalización mediante la reducción de la vista de la red para el ámbito local. Con AMP, un nodo de cualquier red X no sabe sobre el estado de todos los enlaces de la red, pero solamente es informado de sus nodos vecinos. En este sentido, se basa en la propagación de la información de la congestión en la red en el llamado mecanismo de contrapresión (Ver Gráfica 31)

Gráfica 31

Idea Básica del Mecanismo Contrapresión



Tomada de Artículo: Towards low-complexity Internet Traffic Engineering The Adaptive Multi Path Algorithm. De: ScienceDirect 2008. Pág.5

Haciendo una analogía de este concepto se considera una red IP como un sistema unidireccional de malla de tubos poroso que transporta fluido viscoso. Tan pronto como el líquido presenta resistencia, por ejemplo debido a una sección

reducida o a un obstáculo (que corresponde a una situación de congestión en la red) se crea instantáneamente una presión local con dos consecuencias: la presión comienza a propagarse en dirección opuesta al flujo (contrapresión). En segundo lugar si la presión alta persistente (congestión) puede superar la viscosidad del fluido y hacer que se escape a nivel local a través de los poros lo cual corresponde a la pérdida de paquetes en la red. Finalmente ambos efectos dan lugar a la creación de una nueva presión global / pérdida de equilibrio.

Consideremos que el nodo X cuenta con un conjunto Ω_x de nodos vecinos $Y_i \in \Omega_x$, $i \geq 0$, y dejar que $X Y_i$ denote el enlace dirigido genérico del nodo X al nodo Y_i (Ver Gráfica 31). Considerando que con OMP al aumentar la utilización del enlace $Y_0 X$ se produce un impacto en dicho enlace, de tal forma que con AMP el único nodo a reaccionar de forma inmediata es Y_0 como nodo final de $Y_0 X$, tratando de desviar el tráfico de $Y_0 X$ por otros caminos alternos. Adicionalmente el nodo Y_0 envía periódicamente los mensajes llamados contrapresión (BMs- backpressure messages) para cada nodo adyacente $N_i \in \Omega_{Y_0}$ a fin de informar a N_i acerca de su contribución a la situación de congestión en el enlace $Y_0 X$. Todo $N_i \in \Omega_{Y_0}$ a su vez transmiten esta información a sus nodos vecinos, con el fin de poder contribuir en la solución de la congestión presentada.

De esta forma AMP aplica ingeniería de tráfico de forma autónoma en cada nodo de la red, manteniendo el intercambio de información de carga de manera local permitiendo así una baja señalización en la red. Sin embargo el algoritmo no pone limitación alguna sobre el tipo de tráfico de entrada, tolerando matrices de tráfico con un grado arbitrario de variabilidad en los paquetes transmitidos.

Es importante conocer cada una de las partes del algoritmo de enrutamiento adaptativo multicamino AMP, incluyendo el cálculo de múltiples caminos, métricas de carga del enlace, mecanismo de señalización, reenvío de paquetes y el balanceo de carga.

8.4. Cálculo de múltiples caminos

Uno de los requisitos fundamentales de cualquier esquema de ingeniería de tráfico es que la topología de la red muestre una cantidad suficiente de redundancia que permita

desviar tráfico hacia caminos alternativos, en caso de que uno o varios enlaces de la red se congestionen. Sin embargo, la redundancia topológica no es suficiente por sí sola como arquitectura de enrutamiento sino que es necesario aplicar e incorporar principios y mecanismos que lo hagan posible.

Uno de los más importantes objetivos al desarrollar AMP fue permanecer dentro de la actual arquitectura de enrutamiento intradominio y comparar con ECMP el aumento de manera cuidadosa del número de trayectorias múltiples. En este contexto, es importante notar que a diferencia de los enfoques orientados a la conexión como Multi-Protocol Label Switching (MPLS), estos imponen severas limitaciones con respecto al número de opciones multi-ruta. Es decir con MPLS todas las posibles combinaciones de caminos son viables en una red en particular. AMP ofrece la opción de utilizar el mejor camino encontrado. Este criterio establece que en cada nodo de la red cada nodo vecino que está más cerca del nodo destino (en términos de la suma de pesos de enlace) puede ser utilizado como un próximo salto viable para reenviar el tráfico hacia ese destino. Aparte de su gran potencial para aumentar el número de rutas en la red, este criterio garantiza que se eviten los bucles de enrutamiento y que la distancia de cada paquete a su destino siempre disminuya a lo largo de cada salto de su ruta.

8.5. Métricas de carga de enlace – Cálculo de la carga en el enlace

Seleccionar una buena métrica para la carga del enlace es crucial para la toma de decisiones en relación al equilibrio de carga. Esto ocurre especialmente en el tráfico elástico TCP, donde todas las conexiones tienden a aprovechar al máximo la capacidad disponible, por tanto la utilización simple del enlace está dada por:

$$\rho = \frac{\text{Volumen_de_tráfico_transportado}}{\text{Capacidad_del_enlace} * \text{Periodo_de_tiempo}}$$

(Ecuación 18).

A menudo dicha capacidad se aproxima al límite que es el 100%. Sin embargo si dos o más enlaces son iguales, es decir llegan cerca al 100% de la utilización, esto no necesariamente significa que sean la misma carga, (los niveles de carga ofrecida por las diferentes fuentes de tráfico pueden variar considerablemente) y el algoritmo de balanceo de carga debe detectar esto correctamente. Afortunadamente, en el caso de tráfico TCP que es predominante en la Internet de hoy, los niveles de carga ofrecida pueden ser estimados de manera eficiente mediante la evaluación de las probabilidades de pérdida de paquetes en los enlaces, basada en el hecho de que en la medida que exista congestión TCP, la transmisión se retrasa aproximadamente en proporción a la raíz cuadrada de la probabilidad de pérdida de paquetes observados en el enlace. Para los niveles de pérdida de paquetes inferiores al 1%, se supone que TCP no introduce ninguna

desaceleración, lo que significa que por debajo de este límite de utilización simple del enlace la ecuación 18 resulta ser suficiente para estimar la carga ofrecida por el enlace. Resumiendo estas observaciones de la dinámica de TCP, la siguiente métrica para estimar la carga del enlace denominada *carga equivalente* ρ' está dada por:

$$\rho' = \max \left\{ \rho, \rho * K * \sqrt{P} \right\} \quad (\text{Ecuación 19}).$$

En esta ecuación el factor de escala K determina el límite de pérdida de paquetes en donde ρ' excede la utilización total del enlace, por tanto:

$$K > 1/\sqrt{P} \quad (\text{Ecuación 20}).$$

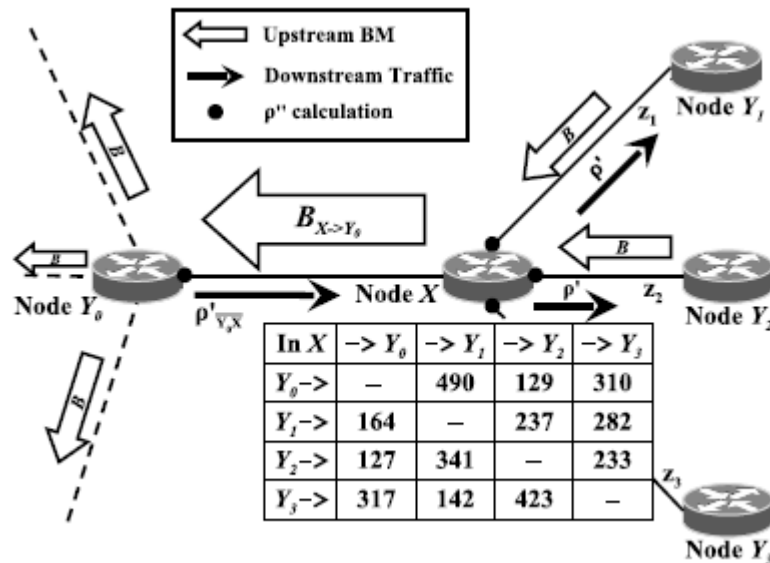
De esta forma se observa que la métrica es muy sensible a las pérdidas de paquetes en el enlace.

8.6. Señalización para enrutamiento adaptativo multicamino (AMP)

Se considera el enlace Y_0X de la gráfica 32 como un ejemplo para describir la información requerida por Y_0 para cada uno de sus enlaces de salida. Y_0 requiere esencialmente dos tipos de información: primero la carga equivalente sobre el enlace Y_0X y segundo la información acerca de la medida del tráfico enrutado desde Y_0 vía X el cual contribuye a la congestión de los enlaces XY_i , $i > 0$. Como Y_0 puede medir y calcular directamente ρ' para su propio enlace Y_0X , el resto de información requerida por Y_0 tiene que ser obtenida de otros nodos por medio de la señalización. Para ello, AMP introduce un tipo específico de señalización de mensajes enviados entre los nodos adyacentes llamados backpressure messages (mensajes contrapresión-BMs). El mensaje contrapresión enviado desde el nodo X hasta el nodo Y_0 debe contener tanto la información directamente relacionada sobre la situación de carga explícita de los enlaces XY_i , $i \geq 1$, así como la información indirecta obtenida por el nodo X sobre la situación de los enlaces lejanos informada por los respectivos nodos Y_i , $i=1,2,3,\dots,n$, donde n es el número de enlaces de salida para el nodo X.

Gráfica 32

Ejemplo para un mensaje contrapresión (BMs- backpressure messages) desde X hasta Y_0 y Matriz de entradas y salidas del nodo X.



Tomada de Artículo: Towards low-complexity Internet Traffic Engineering The Adaptive Multi Path Algorithm. De: ScienceDirect 2008. Pág.7

8.7. Reenvío de paquetes

Un mecanismo apropiado de reenvío de paquetes es necesario con el fin de distribuir el tráfico en las posibles rutas. La técnica de reenvío de paquetes llamada Round Robin representa el mecanismo más simple, pero solo es aplicable si las diferencias de demora en los caminos son muy pequeñas, de lo contrario las conexiones TCP experimentarán desorden en los paquetes y el desempeño será muy pobre.

Existe una segunda técnica que consiste en la división de los prefijos de destino entre múltiples saltos; esta no es muy recomendada debido a que prefijos de destinos IP muy cortos llevará a una granularidad muy gruesa para el desplazamiento de la carga.

El tercer mecanismo consiste en dividir el tráfico de acuerdo a una función hash aplicada entre la fuente y el destino. Este ofrece propiedades muy deseables para

la distribución del tráfico. La función más utilizada es la llamada 16 Bit Cyclic Redundancy Check (CRC-16), la cual permite un tráfico IP realista haciendo que esta sea la primera elección para el equilibrio de carga en Internet. Uno de los más importantes beneficios de esta técnica es su estructura de 16 bits, que permite granularidad muy alta de desplazamiento de carga como $1 / 2^{16} = 1 / 65536$ fracciones de tráfico que puede ser manipulado.

8.8. Balanceo de carga (AMP)

Considerando que la arquitectura de señalización AMP representa un enfoque totalmente nuevo y original, con el propósito de equilibrar la carga se aplica una versión modificada del algoritmo de OSPF - Optimized Multipath (OSPF-OMP), en el cual las decisiones de balanceo de carga se basan exclusivamente en la visión local de un nodo de la red, mientras que AMP logra equilibrar la carga a partir de sus respectivas mediciones y en base a los mensajes contrapresión (Bm's) vistos anteriormente. (Ver Gráfica 31). La motivación principal para recurrir a un algoritmo existente radica en las siguientes ventajas: el algoritmo propuesto logra tanto la convergencia rápida hacia un punto de revisión en la distribución del tráfico, y también opera de forma estable, incluso en la presencia de patrones de tráfico muy ruidoso. Por tanto el algoritmo de equilibrio de carga de AMP logra igualar el valor efectivo de carga equivalente para todos los enlaces de salida.

Siempre AMP debe reajustar las cuotas de tráfico para un destino concreto, el nodo realiza los ajustes que van más allá de considerar solamente la distribución del tráfico actual y la información de señalización. Por otra parte, el nodo realiza un seguimiento de las acciones anteriores con el fin de asegurar una rápida pero estable convergencia hacia un equilibrio de carga en cada enlace de salida; siempre y cuando los enlaces de salida en cada nodo tengan espacio para aceptar tráfico crítico o tráfico de más. El porcentaje del tráfico total desplazado entre los enlaces con menos carga en un periodo corto crece de forma exponencial. En caso de recibir cambios en el estado de los enlaces, los cuales resultan ser críticos el algoritmo de balanceo de carga iniciará la transferencia de tráfico sobre los otros enlaces. Finalmente cuando el siguiente enlace se vuelve crítico otra vez se realizan cambios en la dirección de desplazamiento de carga, etc.

En cuanto a la implementación del algoritmo se refiere, se puede resaltar que está basado en el almacenamiento del balanceo de carga actual para cada enlace de salida hacia un destino específico.

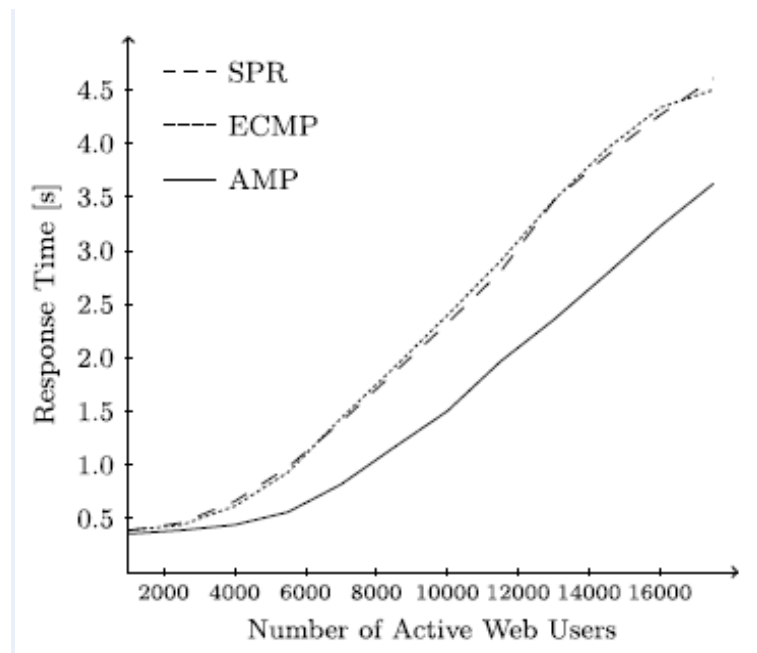
Más allá de posibilitar el crecimiento exponencial de balanceo de carga (y por tanto garantizar la agilidad y la rápida convergencia del algoritmo), la manipulación

de los parámetros de tamaño de paso de carga es también crucial para garantizar la estabilidad del algoritmo y se basa en el siguiente principio: cada vez que un nodo deja de ser crítico, otra vez comienza a aceptar la carga, pero sólo con la mitad de la tasa que había tenido antes de la última vez que llegó a ser crítica. La consecuencia principal de este diseño es que el paso de balancear las cargas disminuirá considerablemente convergiendo hacia la técnica Round Robin cuando el uso de los enlaces resultan ser aproximadamente iguales existiendo poca demora entre ellos.

En la Gráfica 33 se resume los resultados de simulación a nivel de paquetes de datos, que compara el tiempo promedio de respuesta en la web para las tres estrategias de enrutamiento, tales como: Enrutamiento de la ruta más corta (SPR), Equal Cost Multi-Path-Enrutamiento (ECMP) y Adaptive Multi-Path de Enrutamiento (AMP).

Gráfica 33

Promedio de tiempo de respuesta de la página Web de AT & T en EE.UU



Tomada de Artículo: Towards low-complexity Internet Traffic Engineering The Adaptive Multi Path Algorithm. De: ScienceDirect 2008. Pág.10

Se observa en la gráfica anterior que en la fase de simulación inicial (hasta los 2.000 usuarios aproximadamente) el tiempo de respuesta es muy bajo y casi el mismo para las tres técnicas y los resultados restantes corresponden a un tiempo de simulación hasta llegar a los 16.000 usuarios. Las diferentes situaciones en los tiempos de respuesta se caracterizan según el número total de usuarios en el sistema que se muestra en el eje X. El beneficio añadido de AMP en términos del tiempo de respuesta medio para el usuario de la web es sustancial, logrando reducciones de hasta el 43% en comparación con el SPR y ECMP, mientras que las otras dos técnicas de enrutamiento tienen un desempeño casi idéntico.

9. CONCLUSIONES

A lo largo de la historia de las telecomunicaciones, la eficiencia del diseño de la red y la optimización de los recursos empleados han sido el eje de la discusión en la Ciencia. Entre muchas otras cuestiones importantes, el problema de un enrutamiento eficiente en redes basadas en paquetes ha atraído mucho el interés en las últimas décadas. Como resultado, varias propuestas bien fundadas para los protocolos de enrutamiento, incluyendo enrutamiento por vector de distancia y de enrutamiento de estado del enlace, se han desarrollado y aplicado en la práctica. Sin embargo, hasta hace poco, estos protocolos no han tomado en cuenta la distribución real del tráfico, la utilización de los recursos y la congestión en la red. Este hecho ha motivado la aparición de la ingeniería de tráfico como una nueva disciplina que tiene por objeto optimizar el uso de los recursos existentes en las redes operativas.

En el contexto de la ingeniería de tráfico, la idea de protocolos de enrutamiento dinámico es especialmente interesante. Aquí, los mecanismos de distribución de tráfico autónomo dentro de los nodos de la red actúan sobre la información de carga que se intercambia en todo el dominio de la red IP y de esta se optimiza la asignación de recursos en dicha red.

Con el objetivo de combinar las propiedades más favorables de los enfoques existentes, este trabajo describe y evalúa el Adaptive Multi-Path (AMP), como un algoritmo de enrutamiento cuya idea esencial es la de limitar el intercambio de información de carga sólo a los nodos vecinos y reducir así la señalización para el ámbito local. Sin embargo, el resultado de los mensajes de señalización llamados backpressure o de contrapresión también incluyen información sobre el estado de la red más allá del siguiente salto de una manera recursiva, permitiendo así resolver el dilema entre la señalización local y la información de propagación de carga en general en todo el dominio de la red.

Por razones de simplicidad y compatibilidad con la infraestructura de enrutamiento IP actual, el algoritmo Multi-ruta da las especificaciones necesarias para la distribución de la carga mediante el peso de enlaces estáticos.

Además, con el fin de evitar cualquier pérdida de información como reacción a los cambios en los patrones de tráfico, se ha establecido un mecanismo de balanceo de carga para proveer tal situación. El rendimiento de AMP ha sido objeto de una

detallada evaluación técnica de simulación utilizando dos técnicas distintas: simulación a nivel de flujo y a nivel de paquetes.

En el contexto más amplio de redes IP en los próximos años, hay varios temas que serán de gran importancia para toda el área de investigación de ingeniería de tráfico. La construcción de una infraestructura de medición estandarizada y global es ciertamente una de ellas, abriendo el camino hacia el desarrollo de protocolos de enrutamiento de tráfico sensible para redes IP. Sin embargo para monitorear una red estandarizada la solución sería facilitar otras funciones de la ingeniería de tráfico, como la optimización de los pesos de enlaces que resulta ser una de las limitaciones de los sistemas tradicionales actuales y la obtención de una matriz de tráfico de la red en tiempo real.

Las exigencias impuestas a los esquemas de ingeniería de tráfico en el futuro finalmente serán determinadas por los nuevos tipos de aplicaciones que surjan en la red, de su aceptación por los usuarios, y por el desarrollo de tecnologías de conmutación y de transporte más eficientes.

El enrutamiento basado en los pesos de enlaces de las redes IP no ha sido capaz de aprovechar al máximo los múltiples caminos que con frecuencia existen en las redes que proveen servicios de internet. Como resultado, las redes no pueden operar de manera eficiente, especialmente cuando los patrones de tráfico son dinámicos. Ante esto la solución AMP resulta ser favorable para lograr equilibrar la carga de manera global en la red.

En general, la determinación de buenos pesos de enlace a través de diversos métodos ha sido abordada por muchos investigadores. Existen numerosos trabajos en la ingeniería de tráfico IP en los últimos años que consideran diferentes enfoques para el problema de la determinación del peso de los enlaces. La pregunta general es: ¿Se puede determinar un sistema sólido para calcular el peso de los enlaces que trabaje bajo condiciones normales de funcionamiento y también en virtud de una falla o cambio en la matriz de tráfico en la red?

BIBLIOGRAFIA

ANDERSON L, DOOLAN P, FELDMAN L. LDP Specification. 2001. Tomado del link: <http://tools.ietf.org/html/rfc3036>. [Consulta: 12 Noviembre de 2009]

ASH J, GIRISH M, GRAY E. Applicability Statement for CR-LDP. 2002. Tomado del link: <http://tools.ietf.org/html/rfc3213>. [Consulta: 14 Noviembre de 2009]

AWDUCHE D, BERGER L. RSVP-TE: Extensions to RSVP for LSP Tunnels. 2001. Tomado del link: <http://tools.ietf.org/html/rfc3209>. [Consulta: 12 Noviembre de 2009]

BAKER, F. Definición de los Servicios de Campo diferenciada (DS Field) En el IPv4 y IPv6 Cabeceras. 1998. Tomado del link: <http://www.normes-Internet.com/normes.php?rfc=rfc2474&lang=es>. [Consulta: 22 Noviembre de 2009]

BLAKES, BLACK D, CARLSON M. An Architecture for Differentiated Services. 1998. Tomado del link: <http://tools.ietf.org/html/rfc2475>. [Consulta: 27 Noviembre de 2009]

DONOSO M. Yesid. Extensión de los métodos Hop-by-Hop, CR-LDP, y RSVP-Te para Multicast IP sobre MPLS. 2001. Tomado del link: <http://eia.udg.es/~eusebi/papers/clei.pdf>. [Consulta: 29 Noviembre de 2009]

FELICI, Santiago. Evaluación de Mecanismos de Calidad de Servicio en los routers para servicios multimedia. Tomado del link: <http://informatica.uv.es/doctorado/SST/docto-2-qos.ppt>. [Consulta: 25 Noviembre de 2009]

FORTZ Bernard, REXFORD Jennifer y THORUP Mikkel. Artículo: "Engineering With Traditional IP Routing Protocols", Universite Catholique de Louvain, Belgium and AT&T Labs-Research. Florham Park, NJ 07932, Tomado del link: <http://www.sciencedirect.com/> [Consulta: 5 Noviembre de 2009]

GOITIA, María y MARTINEZ, David. “Protocolos de enrutamiento, simulador de tráfico en redes”, Universidad nacional del nordeste, Argentina, Tomado de Internet:

<http://exa.unne.edu.ar/depar/areas/informatica/SistemasOperativos/SimuRedes.pdf> [Consulta: 21 Noviembre de 2009]

GOJMERAC Ivan, REICHL Peter and JASEN Lassen. Artículo: Towards Low-complexity internet traffic engineering the adaptive Multi-Path Algorithm. 2008. Tomado del link:

<http://www.sciencedirect.com/> [Consulta: 15 Enero de 2010]

HAWKINSON, J. Guidelines for Creation, Selection, and Registration of an Autonomous System (AS). 1996. Tomado del link: <http://tools.ietf.org/html/rfc1930>. [Consulta: 20 Noviembre de 2009]

LOBO Castanon. “MPLS”. Universidad Politécnica de Madrid, Departamento de informática – Área ingeniería telemática- 2008, Tomada de Internet: [http://www.it.uniovi.es/docencia/Telematica/com/material/teoria/2008/Tema5 MPLS_telematica.pdf](http://www.it.uniovi.es/docencia/Telematica/com/material/teoria/2008/Tema5_MPLS_telematica.pdf) [Consulta: 11 Noviembre de 2009]

LOZANO, Angélica. TORRES, Vicente y ANTUN Juan Pablo. Tráfico Vehicular. 2003. Tomado del link: <http://www.ejournal.unam.mx/cns/no70/CNS07004.pdf>. [Consulta: 30 Noviembre de 2009]

MALKIN, G. RIP Version 2.1998. Tomado del link: <http://tools.ietf.org/html/rfc2453>. [Consulta: 15 Noviembre de 2009]

MANKIN, A. BAKER, F. BRADEN B. Resource ReSerVation Protocol. 1997. Tomado del link: <http://tools.ietf.org/html/rfc2208>. [Consulta: 10 Noviembre de 2009]

MEDHI Deepankar y RAMASAMY Karthikeyan, “Network Routing. Algorithms, Protocols and Architectures” Libro de la serie: Morgan Kaufmann Publications. 2007. Tomado de internet:

http://books.google.com.co/books?id=IM-Y2W0RIF0C&dq=Network+Routing.+Algorithms.+Protocols+and+Architectures&printsec=frontcover&source=bn&hl=es&ei=OFCeS-iyIsP-8Ab0vKG7Cg&sa=X&oi=book_result&ct=result&resnum=5&ved=0CBwQ6AEwBA#v=onepage&q=&f=false [Consulta: 4 Diciembre 2009]

MILLS, D. Exterior Gateway Protocol Formal Specification. 1984.
Tomado del link: <http://tools.ietf.org/html/rfc904>. [Consulta: 20 Noviembre de 2009]

MOY, J. OSPF Version 2. 1998. Tomado del link: <http://tools.ietf.org/html/rfc2328>. [Consulta: 16 Noviembre de 2009]

ORAN D. OSI IS-IS Intra-domain Routing Protocol. 1990. Tomado del link: <http://tools.ietf.org/html/rfc1142>. [Consulta: 21 Noviembre de 2009]

PADILLA, Jhon Jairo. Componentes de la Ingeniería de Tráfico.pdf. Universidad Pontificia Bolivariana. Bucaramanga. 2009.

RAMIREZ, Mario, JIMENEZ, Tamara y FERNANDEZ, Natalia, “Protocolos para calidad del servicio”. Centro para el desarrollo de las telecomunicaciones de Castilla y León- CEDETEL, - 2010, Tomada de Internet: http://www.cedotel.es/index.php?option=com_content&view=article&id=144%3Apr-otocolo-para-la-calidad-de-servicio-en-redes-opticas-troncales&catid=26%3Adivulgaciontecnologica&lang=es [Consulta: 2 Diciembre de 2009]

RAMIREZ, Ricardo, “Protocolos de enrutamiento”, Universidad autónoma de Colombia, Solticom, Tomada de internet: www.solticom.com/uts/protocolos.pdf [Consulta: 4 Diciembre de 2009]
REKHTER, Y. Border Gateway Protocol. 1995. Tomado del link: <http://tools.ietf.org/html/rfc1771>. [Consulta: 18 Noviembre de 2009]

ROSENE, Viswanathan A. Multiprotocol Label Switching Architecture.2001.
Tomado del link: <http://tools.ietf.org/html/rfc3031>. [Consulta: 10 Noviembre de 2009]

