



RENDIMIENTO DE CALIDAD DE SERVICIO (QOS) DE IPV6 SOBRE TRÁFICO RTP MEDIANTE EMULACIÓN

Jorge Luis Moreira Calderón
Cristhian Gustavo Minaya Vera
Frank Aquino Cornejo Moreira

Rendimiento de calidad de servicio (QoS) de IPV6 sobre tráfico RTP
mediante emulación

Primera edición, diciembre 2016

Autor

Jorge Luis Moreira Calderón
Cristhian Gustavo Minaya Vera
Frank Aquino Cornejo Moreira

RENDIMIENTO DE CALIDAD DE SERVICIO (QOS) DE IPV6 SOBRE TRÁFICO RTP MEDIANTE EMULACIÓN

Jorge Luis Moreira Calderón
Cristhian Gustavo Minaya Vera
Frank Aquino Cornejo Moreira

PRIMERA EDICIÓN

Libro sometido a revisión de pares académicos.



Edición
Diagramación
Diseño
Publicación

Agradecimiento

Este libro no hubiera sido posible sin la ayuda de Dios por habernos guiado por el camino de la felicidad hasta ahora, a nuestros Padres que con sus esfuerzos nos formaron para que ahora logremos este producto, a nuestras esposas e hijos que también fueron clave principal al momento de tenernos en su momento cierto grado de paciencia a todos ellos por brindarnos su apoyo incondicional en cada momento de nuestras vidas, por ser el motor que nos impulsa a seguir luchando día a día, también se agradece a todas aquellas personas e instituciones que de una u otra forma colaboraron para que este libro sea culminado con éxito.

CAPÍTULO I

Retardo producto de los algoritmos de compresión:	9
Retardo producto de la electrónica de los componentes de red:	13
Jitter - Variaciones en la demora	14
Eco	16
Pérdida de paquetes	17

CAPITULO II

PROTOCOLO DE COMUNICACIÓN IPv6	20
Dirección Local Única (Unique-Local Address)	34
Cabeceras de Extensión	39

CAPITULO III

COMUNICACIÓN	42
--------------	----

CAPITULO IV

IPv4 A IPv6	47
Tunelizado	48

CAPITULO V

IMPLEMENTACIÓN DE CALIDAD DE SERVICIO EN IPv6	57
CÁLCULO DEL NÚMERO DE LLAMADAS CONCURRENTES PARA ADMINISTRACIÓN DEL ANCHO DE BANDA	69
DEFINICIÓN DE LA PLATAFORMA DEL MODELO DE CALIDAD DE SERVICIO, Y TÉCNICA DE QUEUING A UTILIZAR.	72

BIBLIOGRAFÍA 82

PRÓLOGO

Los autores describen desde el análisis teórico los parámetros de calidad de servicio y rendimiento de una red IPv6 en la transmisión de voz mediante la emulación de un modelo de red semejante al del Municipio del Cantón Chone, y bajo un modelo de usuario de la misma entidad. El análisis, se realizó con base en los resultados de tres parámetros críticos para tráfico en tiempo real: Jitter, Delay y Packet Loss. Para la parte operativa, se hizo uso de tecnologías de calidad de servicio adecuadas para personalizar el paso de los paquetes por cada uno de los nodos intermedios de la red, en términos de prioridad, previo a un análisis que permitió determinar tanto el modelo de calidad de servicio, y su respectivo sistema de encolamiento, adecuado para tráfico en tiempo real, como también el respectivo codificador de voz a utilizar, y finalmente el cálculo de ancho de banda para el modelo de usuario planteado. Este análisis se lo realizó tomando como referencia los valores obtenidos de IPv4, para luego compararlos con los de IPv6, y determinar su desempeño.

Cada uno de ellos muestra en este texto un amplio desempeño profesional definiendo la importancia del manejo de la tecnología, la comunicación y la dinámica de trabajo basado en redes de transmisión.

El libro presentar resultados luego de realizar varias pruebas con

ambos protocolos, las mismas mostraron que los protocolos IPv4 e IPv6 tienen similares niveles de rendimientos en un ambiente bien dimensionado en cuanto a ancho de banda y con una selección de tecnologías de calidad de servicio adecuadas para tráfico en tiempo real. Sin embargo, cuando se redimensionó el sistema para el doble de llamadas concurrentes, para lo cual fue construido, los niveles de latencia superaron en un 59,8%, con respecto a IPv4. En el caso de IPv6, se redujo casi tres veces la pérdida de paquetes, con respecto a los obtenidos con IPv4.

CAPÍTULO I

Fenómenos que afectan el desempeño de la transmisión de paquetes

Este capítulo permite identificar el estado teórico del desempeño de transmisión, definiendo los conceptos que enmarcan la investigación. para el caso de análisis de VoIP y Telefonía IP, se ha determinado que estos se ven sumamente perjudicados por Retardo (delay), Variaciones en la demora (Jitter), Eco y Pérdida de paquetes (Packet Loss). A continuación, se definen los conceptos e implicaciones de estos tres fenómenos.

Al hablar de transmisión de voz sobre redes Ethernet, el retardo o demora en la llegada de paquetes es un factor sobresaliente para definir la calidad en las comunicaciones. De manera global el retardo está determinado por algunos factores, a saber:

Retardo producto de los algoritmos de compresión:

Las comunicaciones de voz, en su forma natural, son un conjunto de datos analógicos con un sinnúmero de variaciones que, al llegar al cerebro humano, cobran sentido completo. El reto se focaliza en llevar esos datos analógicos a través de una red digital (Ethernet) sin que se pierda o distorsione la información contenida desde su origen. Para lograr el cometido se han desarrollado diversos algoritmos de

compresión que actúan como codificador-decodificador (CODEC). Los CODEC además de convertir datos analógicos en digitales, realizan la compresión de las secuencias de datos, lo cual genera un ahorro en el ancho de banda utilizado, y proporcionan la cancelación del eco. La característica de compresión de datos realizada por los CODEC, ayuda a las redes que poseen enlaces de baja capacidad a mantener más y mejores conexiones VoIP de forma simultánea. Otra forma de minimizar el uso de ancho de banda por las comunicaciones de voz, es evitar el envío de información cuando se producen silencios en medio de las comunicaciones humanas.

Para generar una idea más clara de lo estipulado en la tabla anterior, se definen a continuación algunos conceptos:

Bit Rate (Kbps): Indica la cantidad de bits (información) que se envía por segundo.

Sampling Rate (KHz): Señala la frecuencia de muestreo de la señal de voz, es decir cada cuanto tiempo se toma una muestra de la señal analógica original.

Frame size (ms): Indica cada cuánto tiempo, medido en milisegundos, se envía un paquete con la información digitalizada de la señal de voz.

Mean Opinion Score (MOS): Es un sistema que clasifica la calidad de la voz de las conexiones telefónicas. Con MOS, una amplia gama de oyentes juzga la calidad de una muestra de voz mediante una escala que va del 1 (mala) al 5 (excelente). Los puntajes se promedian para brindar una MOS para el CODEC.

Mientras más compresión de la data realice el CODEC, la latencia introducida será mayor puesto que el tiempo de procesamiento es

más alto.

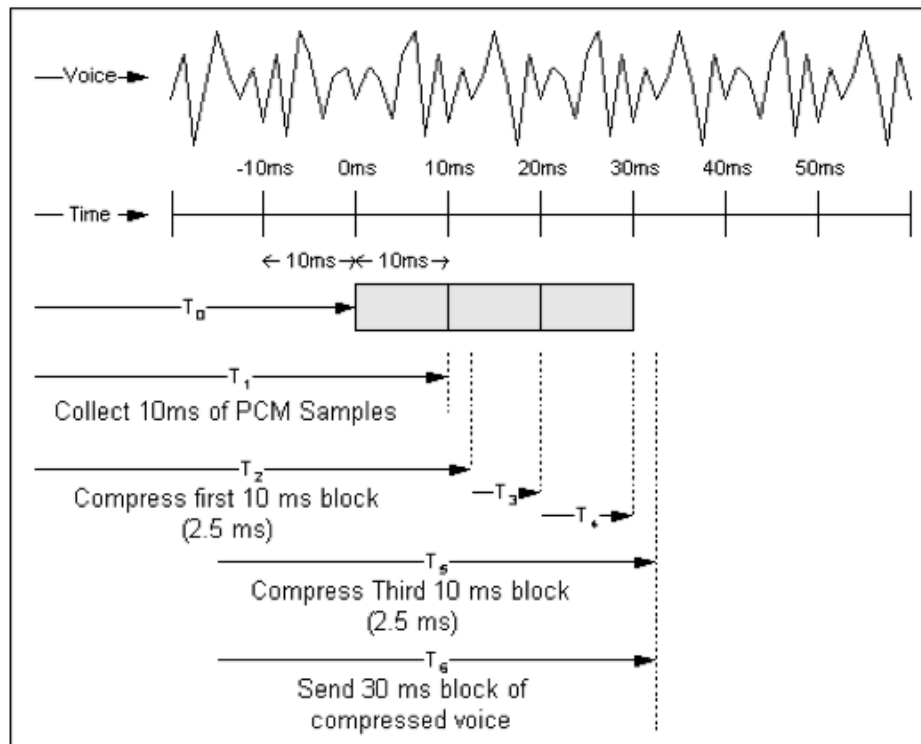
Retardo por empaquetamiento:

Corresponde al tiempo empleado para llenar un paquete de información, denominado carga útil (payload), de la conversación ya codificada y comprimida. Este delay es función del tamaño de bloque requerido por el codificador de voz y el número de bloques de una sola trama. Dentro del protocolo de tiempo real, RTP por sus siglas en inglés, las muestras de voz con frecuencia son acumuladas antes de ponerlas en una trama para su transmisión; esto con el propósito de reducir la cantidad de cabeceras (overhead). La RFC 1890 especifica que el retardo de empaquetamiento por defecto debiera ser de 20 ms. Para G.711, esto significa que 160 muestras serán acumuladas y solo entonces transmitidas en una sola trama.

Se debe mencionar que, al pasar de voz analógica a voz digital, todas las muestras sufrirán tanto de retardos debido al CODEC cuanto al empaquetamiento, pero, que en realidad, estos efectos se superponen en la línea del tiempo. Esto se verifica en el siguiente gráfico:

Gráfico 1: Acumulación de retardos CODEC + Empaquetamiento

Cisco (2006, párr. 19).



Retardo de Serialización:

Corresponde al tiempo requerido para transmitir un paquete IP y está vinculado directamente con la tasa del reloj de la transmisión.

El retardo de serialización se presenta cuando los paquetes pasan a través de diversos dispositivos de almacenamiento y retransmisión, como un Router o un Switch. Así, una trama que atraviesa 10 equipos interconectados, incurrirá en 10 veces el retardo teórico de un solo equipo.

Retardo de Propagación:

Se define como el tiempo que necesita la señal digital para viajar a lo largo de un medio de transmisión. Se encuentra estrechamente relacionado con la distancia geográfica. La velocidad de propagación en el cable de cobre es aproximadamente de 4 a 6 ms/Km.

Retardo producto de la electrónica de los componentes de red:

Son retardos causados por la propia electrónica del equipo de red (NIC, router, switch etc.) al momento de efectuar el proceso de transmisión / recepción. Por ejemplo, una trama que está pasando a través de un router tiene que moverlo desde el puerto de entrada al puerto de salida a través del backplane. Existe un retardo mínimo en el paso por el backplane y algunos retardos variables debidos al encolamiento y procesamiento en el router.

El límite generalmente aceptado para retardos en las conexiones de voz es de 150ms en un solo sentido (o 250ms como límite máximo) para que se las siga considerando de buena calidad. Cuando los retardos se incrementan más allá de estos límites, los hablantes y los escuchas empiezan a de-sincronizarse, y muchas veces empiezan a hablar al mismo tiempo, o ambos esperan a que el otro hable. Esta condición es comúnmente llamada solapamiento del hablante (solapamiento de la conversación). Este efecto puede ser observado en llamadas telefónicas internacionales las cuales viajan a través de conexiones satelitales. Es necesario mencionar que los retardos en este tipo de enlaces satelitales se encuentran en el orden de 500ms, 250ms de subida y 250ms en el retorno.

De forma general, la ITU (International Telecommunications Union) ha definido retardos aceptables a aquellos que no superan los 150 ms en un solo sentido. Cisco (2006).

Jitter - Variaciones en la demora

Se denomina jitter a la variación en las latencias. El jitter se define técnicamente como la variación de tiempo en la llegada de paquetes, causada por congestión en la red, pérdida de sincronización o por las diferentes rutas seguidas por los paquetes para llegar al destino (Elastix, s.f). Teóricamente un paquete debe salir del origen cada cierta cantidad de tiempo, y llegar al destino en periodos constantes iguales a los del origen; sin embargo, esto no ocurre en la realidad. A esta diferencia entre las variaciones, se denomina jitter. Este fenómeno se ocasiona debido a la cantidad de retardo en las colas y el tiempo de procesamiento, que pueden variar en función del tráfico que corre en la red. Ejemplo: Si un Gateway de voz transmite tramas a intervalos constantes (20 ms), el Gateway destino, de manera general, no recibirá esas tramas a intervalos regulares debido al problema de jitter.

En cuanto a los valores recomendados para jitter existen diversos conceptos, algunos más estrictos que otros, que estipulan los valores máximos que debería alcanzar este parámetro en la red para soportar el servicio de voz (VoIP o Telefonía IP).

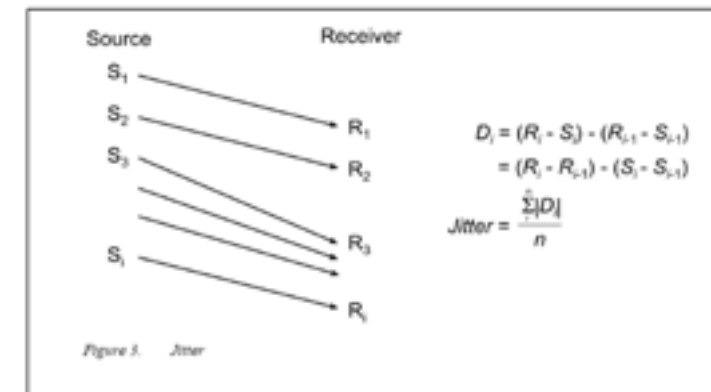
Especialistas en centralitas IP tales como 3CX (s.f) definen que “El jitter entre el punto inicial y final de la comunicación debe ser inferior a 100 ms. Si el valor es menor a 100ms el jitter puede ser compensado

de manera apropiada. En caso contrario debería ser minimizado” .

Esto se ilustra en la figura siguiente:

Gráfico 2: Definición de Jitter

(Rosario, M. A., 2006, cap. 3)



Con el propósito de minorar el problema de jitter, se almacena las tramas recibidas en un buffer lo suficientemente grande que permita recibir las tramas más lentas y ubicarlas de manera ordenada. Para minimizar el efecto de retardo debido al almacenamiento en el buffer, algunas aplicaciones utilizan un sistema adaptativo. Es decir, si el conjunto de jitter en la red es pequeño, el tamaño del buffer será también pequeño, pero si, debido al aumento de tráfico en la red, el jitter empieza a elevarse, el buffer de destino será incrementado de forma automática con el fin de compensar ese aumento. En tal sentido, el jitter en la red empeorará la calidad de la comunicación de voz en la misma proporción que crece el retardo de extremo a extremo debido al buffering en el destino.

Eco

El efecto Eco dentro de una llamada telefónica de cualquier índole, sea esta analógica, IP, satelital etc., es producto del retardo dentro de las redes de comunicaciones. Existen algunos tipos de eco, entre ellos citaremos:

Eco acústico: Producido generalmente por acoplamientos inadecuados entre los dispositivos de habla y escucha.

Eco híbrido: El que se presenta debido a residuos de energía eléctrica que regresan hacia el circuito híbrido en la PSTN.

Cuando los retardos en la red son cortos (menores a 25ms), el eco generado por el circuito de voz regresará muy rápidamente al generador de la llamada y no será percibido. Sin embargo, dentro de una red VoIP, la media de los retardos siempre supera los 25 ms por lo cual es necesario disponer de un método de cancelación de eco. Como habíamos dicho con anterioridad, los CODEC también realizan el trabajo de cancelación de eco.

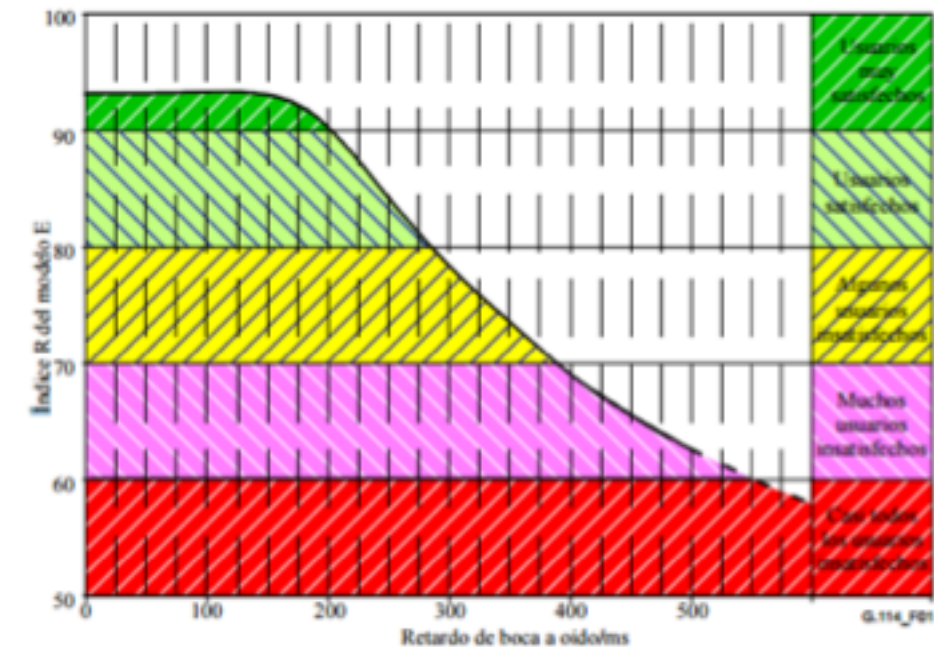
Incluso con un excelente método de cancelación, puede presentarse otro fenómeno dentro de las comunicaciones de voz, este es el solapamiento de la conversación (talker overlap).

Este problema se da cuando las conversaciones de los abonados se superponen debido al gran retardo existente en la red. La recomendación de la ITU G.114: Tiempo de transmisión en un sentido, define que “independientemente del tipo de aplicación, se recomienda que el retardo en un sentido no supere 400 ms para la planificación general de la red”.

En lo que respecta a las aplicaciones de voz, el siguiente gráfico

muestra la relación entre la satisfacción del usuario y los retardos de un solo sentido en la red.

Gráfico 3: Satisfacción del usuario vs. Retardos en un solo sentido
UIT-T G.114 (2004, p. 3).



Pérdida de paquetes

En las redes de telefonía PSTN los canales de comunicación son bastante seguros y estables, lo cual garantiza que las conversaciones analógicas se lleven a cabo con un nivel de confiabilidad que bordea el 100%.

Por su parte las redes de datos admiten dentro de su concepción que existan paquetes que se pierdan al momento de realizar el proceso de transmisión / recepción. De manera general para Ethernet, cuando los paquetes son perdidos, los protocolos de alto nivel aseguran que estos sean reenviados; sin embargo, dentro de los protocolos

diseñados para aplicaciones RTP, el proceso de confirmación de que un paquete llegó o no es inexistente, porque el tráfico de confirmación carga inútilmente a la red de datos, y además, se considera que el canal fue pensado para ser lo suficientemente seguro con el propósito de soportar comunicaciones en tiempo real.

En las aplicaciones de tiempo real como voz y video, la data es encapsulada en paquetes y enviados sin que exista la confirmación de recepción. Cuando el porcentaje de pérdida de paquetes es pequeña, la degradación de la data es imperceptible.

Sin embargo, hay que ser muy claros en que VoIP y Telefonía IP son intolerantes a la pérdida de paquetes. Incluso un 1% de pérdida de paquetes puede degradar de manera significativa una llamada de voz usando el códec G.711. Otros códec que emplean más compresión pueden tolerar incluso menos pérdida. Según Cisco: “El códec G.729 por defecto requiere una pérdida de paquetes mucho menor al 1% para evitar errores audibles. Idealmente, no debería existir pérdidas de paquetes para VoIP” (VOIP Wiki. 2016).

El efecto de packet loss se debe principalmente a una red congestionada o al exceso de ruido eléctrico dentro de los dispositivos o cableado de red.

En el ámbito de las comunicaciones de voz, existen diversos mecanismos que evitan la degradación de las comunicaciones producto de la pérdida de paquetes. Una de las más comunes es la retransmisión del último paquete recibido.

Finalmente, cabe mencionar que se cuentan como paquetes perdidos las tramas que llegan a destiempo o fuera de orden.

CAPÍTULO II

PROTOCOLO DE COMUNICACIÓN IPv6

El protocolo de Internet versión 6 (IPv6) es la última versión de los protocolos IP y la primera versión de protocolos que es ampliamente desplegado. IPv6 fue desarrollado por el IETF (Internet Engineering Task Force) para hacer frente al problema de agotamiento de direcciones IPv4.

El protocolo de Internet versión 6 es un nuevo protocolo de direcciones diseñado para incorporar todos los posibles requisitos de la futura Internet conocido en el mundo de IT como versión de Internet 2. Este protocolo, como su predecesor el IPv4, trabaja en la capa de red (capa 3). Junto con su oferta de una enorme cantidad de espacio de direccionamiento lógico, este protocolo tiene amplias características que se ocupan de las deficiencias que tiene IPv4.

Hasta la fecha, el protocolo de Internet ha sido reconocido únicamente como IPv4. Las versiones de la 0 a la 3 fueron usadas mientras el protocolo estaba siendo desarrollado y experimentado. Por tanto, podemos asumir gran cantidad de actividades detrás del posicionamiento de un protocolo en ambiente de producción. De forma similar, el protocolo IP versión 5 fue usado mientras se experimentaba con el corriente protocolo de Internet.

Después del desarrollo de IPv4 en los inicios de 1980, la cantidad disponible de direcciones IPv4 empezó a reducirse rápidamente, así como la demanda de direcciones creció exponencialmente con la divulgación masiva de la Internet. Tomando en cuenta esta situación, la IETF, en 1994, inició el desarrollo de un protocolo de

direccionamientos que reemplace a IPv4. El progreso de IPv6 puede ser seguido en las publicaciones RFC (Request for Comments) de la IETF:

- 1998 – RFC 2460 – Basic Protocol (Protocolo Básico)
- 2003 – RFC 2553 – Basic Socket API (Socket API Básico)
- 2003 – RFC 3315 – DHCPv6 (DHCP versión 6)
- 2004 – RFC 3775 – Mobile IPv6 (IPv6 Móvil)
- 2004 – RFC 3697 – Flow Label Specification (Especificación de etiquetado de flujo)
- 2006 – RFC 4291 – Address architecture (Arquitectura de direccionamiento)
- 2006 – RFC 4294 – Node requirement (Requerimiento de nodo)

El 6 de junio de 2012, algunos de los gigantes de la Internet (Google, AT&T y CISCO) escogieron colocar sus servidores en IPv6. Course Hero (2015). A la presente están utilizando el mecanismo Dual Stack (Doble Apilamiento) para implementar IPv6 en paralelo con IPv4.

La característica de IPV6 es que mejora la comunicación en tiempo real al ser sucesor de IP no ha sido diseñado para ser compatible hacia atrás. Intentando mantener las funcionalidades básicas del direccionamiento IP, el nuevo protocolo IPv6 fue rediseñado íntegramente. IPv6 ofrece las siguientes características que mejoran la comunicación de las tramas RTP:

No Fragmentation – Ausencia de Fragmentación

En IPv6, los routers (internos o de núcleo) no efectúan

fragmentación de paquetes, en su lugar, la fragmentación es realizada de extremo a extremo. Esto quiere decir que los nodos fuente y destino realizan, a través de la pila de IPv6, la fragmentación de un paquete y luego el re-ensamble, respectivamente. El proceso de fragmentación consiste en dividir en paquetes más pequeños la parte “fragmentable” del paquete fuente, y adicionar a cada una la parte “no fragmentable”.

IPv6 necesita que cada enlace en la Internet tenga un MTU (Maximum Transmission Unit) de 1280 octetos o mayor. Sobre cualquier enlace que no pueda transmitir un paquete de 1280 octetos en una sola pieza, se debe proveer, bajo la capa de IPv6, la fragmentación y re-ensamble del enlace específico.

Enlaces que tienen una Unidad Máxima de Transmisión (MTU) configurable, por ejemplo, enlaces PPP definidos en la RFC 1661, deben ser configurados para tener una MTU de al menos 1280 octetos, aunque es recomendable que estos sean configurados con una MTU de 1500 octetos o mayores, para dar cabida a posibles encapsulamientos (Ejemplo: tunelizado) sin incurrir en la fragmentación de la capa de IPv6. Desde cada enlace al cual el nodo está directamente ligado, el nodo debe ser capaz de aceptar paquetes tan extensos como el MTU del enlace.

Es extremadamente recomendable que los nodos IPv6 implementen Path MTU Discovery (Descubrimiento de Ruta de MTU – RFC1981), con el propósito de descubrir y tomar ventaja de las rutas con MTU mayores a 1280 octetos. Sin embargo, una mínima implementación de IPv6 (en una boot ROM) puede simplemente restringirse a sí misma para enviar paquetes no mayores de 1280 octetos, y omitir la implementación de Path MTU Discovery.

Para enviar un paquete mayor que el MTU de la ruta, un nodo puede utilizar el encabezado Fragment para fragmentar el paquete en la fuente y tenerlo desfragmentado en el o los destinos. Sin embargo, el uso de dicha fragmentación es evitado en cualquier aplicación que es capaz de ajustar sus paquetes acorde a la medida del MTU de la ruta.

Un nodo debe ser capaz de aceptar paquetes fragmentados que, después de desfragmentarse, sean tan grandes como 1500 octetos. A un nodo le está permitido aceptar paquetes fragmentados que re-ensamblen más de 1500 octetos. Un protocolo de capa superior, o aplicación, que dependa de la fragmentación de IPv6 para enviar paquetes mayores al MTU de la ruta no deberá enviar paquetes mayores a 1500 octetos a menos que sea seguro que el destino es capaz de re-armar paquetes de ese tamaño.

En respuesta a un paquete IPv6 que es enviado a un destino IPv4, como por ejemplo un paquete que experimenta una traducción de IPv6 a IPv4, el nodo IPv6 origen puede recibir un mensaje ICMP (Internet Control Message Protocol) denominado Packet Too Big (Paquete demasiado grande) reportando un Next-Hop MTU menor de 1280. En ese caso, el nodo IPv6 no requiere reducir el tamaño de los paquetes subsiguientes a menos de 1280, pero debe incluir un encabezado de fragmento en esos paquetes para que el router que hace la traducción de IPv6 a IPv4 pueda obtener un valor adecuado de identificación para usarlo en los fragmentos resultantes IPv4. Nótese que esto significa que la carga útil (payload) tiene que ser reducida a 1232 octetos (1280 menos 40 para el encabezado IPv6 y menos 8 para el encabezado fragmento), y aún más pequeño si son utilizados

encabezados de extensión (Minoli. 2006: 288).

Cabecera simplificada

La cabecera del protocolo IPv6 ha sido simplificada removiendo toda la información innecesaria y las opciones, que están presentes en la cabecera de IPv4, al final de la cabecera de IPv6. El encabezado de IPv6 solo llega a ser el doble del tamaño de IPv4 adicionando el hecho de que la dirección IPv6 es 4 veces más grande. Este cambio permite que los routers procesen datagramas de manera más rápida y mejore la velocidad en general, (mlrEluCx, 2009).

Campo checksum eliminado en IPv6: IPv6 mejora los tiempos de latencia. El campo checksum de 16 bits que se utiliza para la verificación de errores del encabezado IP, (CISCO, 2016) ya no aparece en la cabecera IPv6. El motivo principal es la innecesaria redundancia.

En IPv4 se está facilitando la misma información de varias formas, otros mecanismos de encapsulado realizan la misma función, por ejemplo: IEEE 802 MAC, framing PPP, capa de adaptación ATM, etc., (Ortega, 2012) y esto conlleva a menor tiempo de procesamiento en los paquetes.

Conectividad de extremo a extremo

Ahora cada dispositivo en el mundo puede disponer de una dirección IP que puede atravesar la Internet sin utilizar NAT (Network Address Translation – Traducción de direcciones de Red) u otros componentes de traducción. Después de que IPv6 sea implementada

completamente, cada Host podrá directamente llegar hasta otro host en la Internet, con algunas limitaciones como políticas de red, firewall (cortafuegos), etc.

Para proveer mejor soporte al tráfico de tiempo real (Ej.: VoIP, IPTV), IPv6 incluye dentro de sus especificaciones el tráfico etiquetado. A través de este mecanismo, los routers pueden reconocer el tráfico de extremo a extremo al cual pertenecen los paquetes transmitidos. Esto es similar al servicio ofrecido por MPLS, pero esta característica ya está incluida en IPv6 y no como un adicional.

Otro de los problemas actuales tiene relación con la integridad de la señalización y la ruta de la portadora de VoIP, específicamente el hecho de que los paquetes de VoIP tienen dificultades de temporización al ser encaminados a través de corta-fuegos, no solo por las consideraciones del protocolo, sino, a nivel práctico, por problemas referidos a la traducción de direcciones de red (NAT) (Minoli. 2006: 8). Algunos protocolos no pueden viajar a través de dispositivos NAT; dado que, NAT hace que numerosas aplicaciones, como por ejemplo VoIP, no puedan ser efectivamente utilizadas en todas las instancias. Como consecuencia estas aplicaciones pueden ser solamente usadas en el ámbito de la intranet. Algunos ejemplos incluyen:

- Aplicaciones multimedia como video-conferencia, VoIP a través de Internet o video sobre demanda / IPTV no trabajan de manera afinada a través de dispositivos NAT. Las aplicaciones multimedia hacen uso del Protocolo de Tiempo Real (RTP) y el Protocolo de Control de Tiempo Real (RTCP); a su vez, éstos utilizan UDP con asignación dinámica de puertos (NAT no admite directamente este entorno).

- La autenticación necesita la dirección fuente, desafortunadamente, la dirección fuente en la cabecera IP es a menudo modificada por el dispositivo NAT.

- IPSec es utilizada ampliamente para autenticación de datos, integridad y confidencialidad, sin embargo, cuando NAT es utilizado para la operación de IPSec se ve impactada, ya que NAT cambia la dirección en la cabecera IP, por tanto, en la práctica, no es utilizada a menos que sea el caso.

Afortunadamente, en IPv6, NAT desaparece y así también, los inconvenientes que genera.

En definitiva, no hay NAT porque todas las interfaces que acceden a Internet tienen una dirección global asignada, (Gerometta, 2015).

Envío y enrutamiento de mayor velocidad

Como se había mencionado, la cabecera de IPv6 coloca al final de ella la información innecesaria. La información contenida en la primera parte de la cabecera ayuda a los router a tomar decisiones más rápidas en función, únicamente, de la información primaria de la cabecera.

Soporte mejorado de prioridad

En el caso de IPv4, este utiliza 6 bits DSCP (Differential Service Code Point - Código de Servicio Diferencial de Punto) y 2 bits ECN (Explicit Congestion Notification – Notificación Explícita de Congestión) para proporcionar la característica de calidad de servicio, pero sólo podría ser utilizado si los dispositivos de extremo a extremo lo soportan, es decir, los dispositivos origen y destino y la red subyacente deben ser

compatibles con él.

En IPv6, la clase de tráfico y la etiqueta de flujo son utilizadas para mostrar a los routers implicados la forma de procesar de manera eficiente el paquete y cómo enrutarlo. La adición del campo etiqueta de flujo en IPv6, permite identificar comunicaciones que deben ser procesadas de modo particular a lo largo de la ruta, sin necesidad de que los dispositivos intermedios procesen múltiples campos de los encabezados (Gerometta. 2016: 15) (dirección fuente, dirección destino, puerto origen, puerto destino, protocolo), ya que la información se identifica directamente por la etiqueta, lo cual reduce tiempo de procesamiento de los paquetes. La etiqueta de flujo asignada es elegida de forma pseudo aleatoria y uniforme en el rango de 1 a FFFFF hexadecimal (IETF, 1995)

Transición suave

La extensa cantidad de direccionamiento dentro del protocolo IPv6 permite asignar una dirección IP única a cada dispositivo del planeta. Con esta aseveración se admite que los mecanismos como NAT para ahorrar direcciones IP ya no son requeridos. Por lo tanto, los dispositivos pueden enviar/recibir datos desde cualquier otro dispositivo. Por ejemplo, VoIP o cualquier otro flujo de datos RTP pueden ser utilizados mucho más eficientemente.

Otro aspecto es que la cabecera se encuentra menos cargada, así que los routers pueden tomar decisiones de reenvío tan pronto como llegan los paquetes.

MODOS DE DIRECCIONAMIENTO

En las redes de computadoras, el modo de direccionamiento se refiere al mecanismo de cómo se llega a un host en la red. El protocolo IPv6 ofrece algunos tipos de modos con los cuales un único host (unicast) puede ser abordado en la red, más de un host puede ser direccionado a la vez (multicast) o el host ubicado a la distancia más cercana puede ser abordado (anycast).

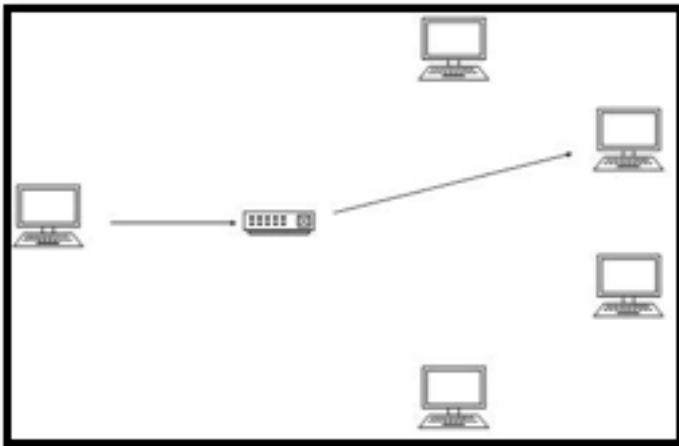
Unicast (unidifusión)

En el modo unicast de direccionamiento, una interface (host) IPv6 se identifica de forma única en un segmento de red. El paquete IPv6 contiene las direcciones IP tanto de la fuente cuanto del destino. La interfaz del host es equipada con una dirección IP la cual es única en el segmento de red. En este ambiente de red, un switch o un router cuando reciben un paquete IP unicast, destinado a un solo host en la red, envían dicho paquete a una de sus interfaces de salida la cual conecta hacia ese host en particular que es el destino.

La siguiente imagen describe el funcionamiento de unicast:

Gráfico 4: Mensajes unicast

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 5



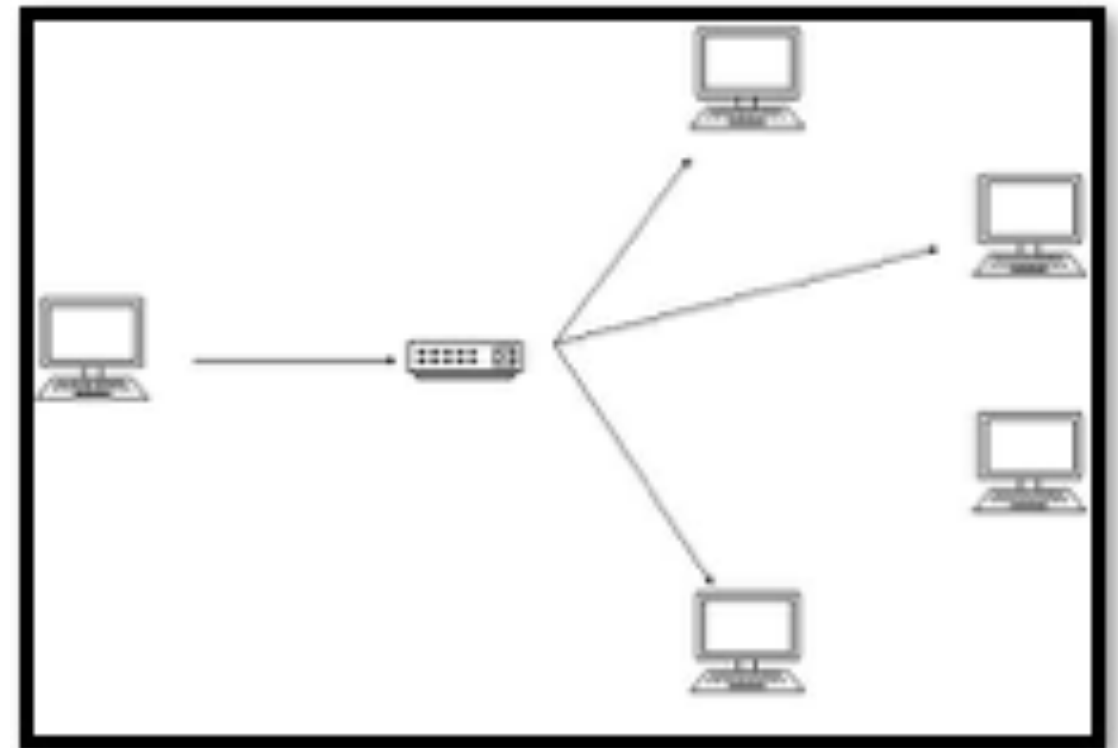
Multicast (multidifusión)

El modelo de multicast en IPv6 es el mismo que en IPv4. El paquete destinado a múltiples hosts es enviado a una dirección especial multicast. Todos los host interesados en esa información multicast, primero deben adherirse al grupo de multicast. Todas las interfaces que se han unido al grupo reciben el paquete multicast y lo procesan, mientras los hosts que no están interesados en los paquetes de multidifusión ignoran dicha información.

A continuación, un gráfico explicativo de multicast:

Gráfico 5: Mensajes multicast

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 6

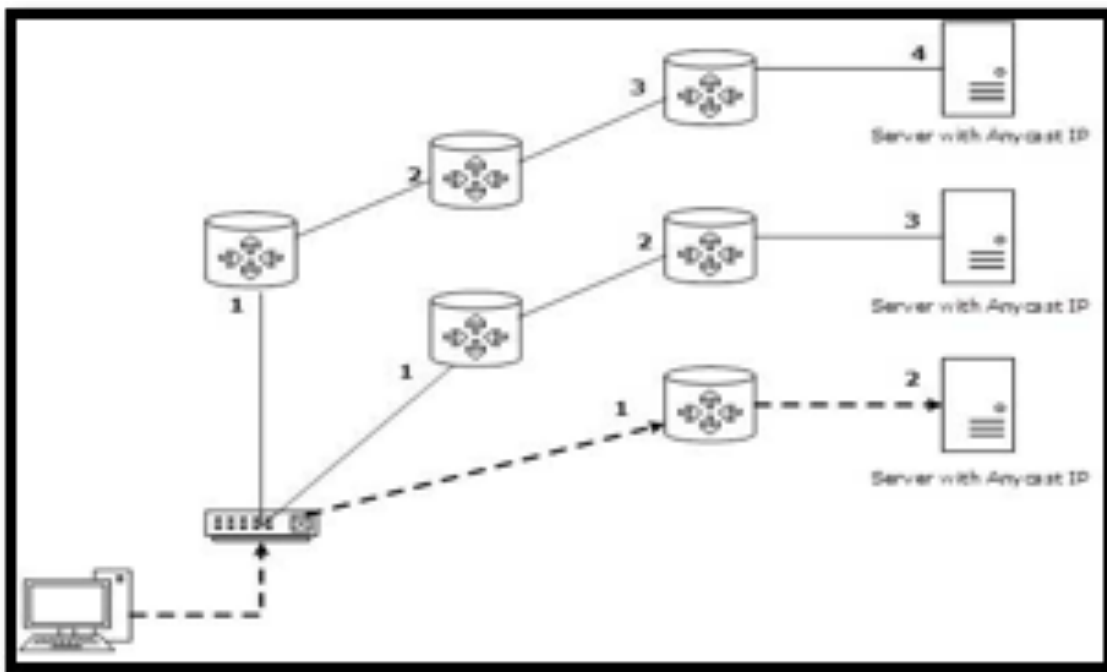


Anycast (cualquier-difusión)

IPv6 ha introducido un nuevo sistema de direccionamiento al cual se lo ha llamado direccionamiento anycast. En este modo de direccionamiento, múltiples interfaces (hosts) han sido asignadas a la misma dirección IP de anycast. Cuando un host desea comunicarse con un host equipado con una dirección IP de anycast, envía un mensaje Unicast. Con la ayuda de un mecanismo complejo de enrutamiento, el mensaje Unicast es entregado al host más cercano al remitente, en términos de costo de enrutamiento.

Gráfico 6: Mensajes anycast

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 6



Un ejemplo de anycast se describe a continuación:

Imaginemos que los servidores Web de tutorialderedes.com se encuentran localizados en todos los continentes. Ahora, asumamos que todos los servidores tienen asignados una sola dirección IPv6 del tipo anycast. Entonces, cuando un usuario de América Central quiera alcanzar tutorialderedes.com, el DNS apuntará hacia el servidor que físicamente se encuentra en América Central. Si un usuario de Europa quisiera acceder al servidor tutorialderedes.com, el DNS apuntaría únicamente al servidor que se encuentra en Europa. Más cercano o en condiciones más próximas son términos que se utilizan en la teoría de costo de ruteo.

En el gráfico 6: Mensajes anycast, cuando un cliente trata de alcanzar al servidor, la solicitud es enviada al servidor con menor costo de ruteo.

SISTEMA DE NUMERACIÓN HEXADECIMAL

Antes de adentrarse en el estudio del formato de las direcciones IPv6, se debe realizar una revisión del sistema de numeración hexadecimal. Este sistema de numeración tiene base 16. Para representar valores en un formato legible, se utilizan símbolos del 1 al 9 para representar valores del uno al nueve y símbolos de la "A" a la "F" para representar valores del diez al quince.

ESTRUCTURA DE LAS DIRECCIONES

Una dirección IPv6 está conformada de 128 bits divididos en ocho bloques de 16 bits. Cada bloque es entonces convertido en un número hexadecimal de 4 dígitos separados por el símbolo dos puntos (:).

Por ejemplo, la siguiente dirección IPv6 de 128 bits es representada en formato binario y dividida en ocho bloques de 16 bits:

Cada bloque es luego convertido en Hexadecimal y separado por el símbolo “:”, así:

2001:0000:3238:DFE1:0063:0000:0000:FEFB

Incluso después de convertir al formato hexadecimal, la dirección IPv6 permanece extensa. IPv6 provee algunas reglas para acortar las direcciones, estas reglas son:

Regla 1: Descartar el o los ceros que quedan a la izquierda

En el bloque número 5 del ejemplo planteado, los dos ceros de la izquierda, pueden ser omitidos, así:

2001:0000:3238:DFE1:63:0000:0000:FEFB

Regla 2: Si dos o más bloques contienen ceros consecutivos, omitirlos a todos y reemplazar con doble signo de “dos puntos”, así los bloques 6 y 7 del ejemplo serían:

2001:0000:3238:DFE1:0063::FEFB

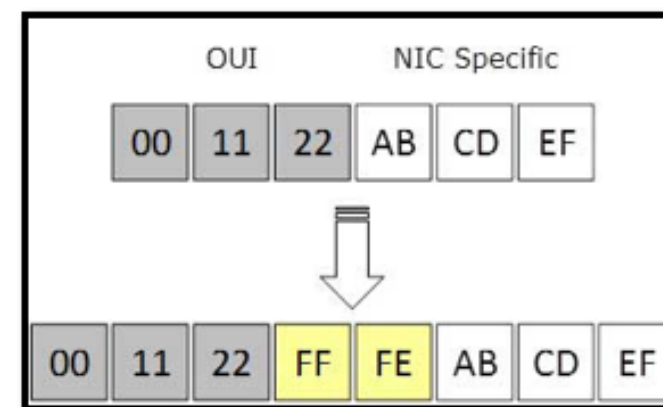
Identificación de Interfaz

IPv6 tiene tres diferentes tipos de esquemas de direccionamiento Unicast (Unidifusión). La segunda mitad de la dirección, es decir los últimos 64 bits, son siempre utilizados para la identificación de la interfaz (Interfaz ID). La dirección MAC de un dispositivo está compuesta de 48 bits y representada en hexadecimal. La dirección MAC se considera asignada como única en todo el mundo, es

decir, dos dispositivos de red no pueden tener la misma dirección MAC. La identificación de interfaz (Interfaz ID) toma ventaja de esta característica. Un host puede auto-configurar su Interfaz ID usando el formato de identificador único extendido (Extended Unique Identifier) EUI-64. Primero, un host divide su propia dirección MAC en dos mitades de 24 bits. Luego el valor hexadecimal 0xFFFE es introducido entre esas mitades de dirección MAC, resultando una Interface ID de 64 bits.

Gráfico 7: EUI-64 Interface ID

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 9

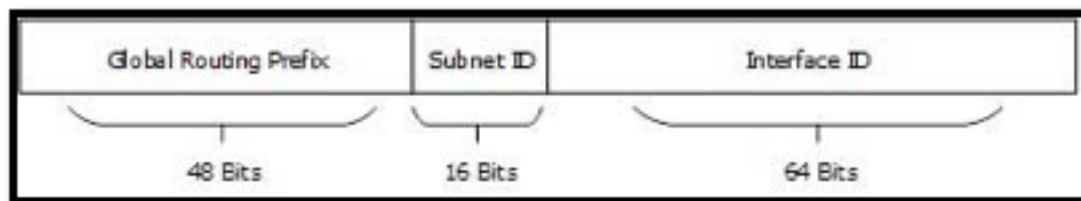


Dirección Unicast Global

Este tipo de dirección es equivalente a una dirección IPv4 pública. En IPv6 las direcciones Unicast Globales (Global Unicast) son mundialmente identificables y direccionables de manera única.

Gráfico 8: Global Unicast Address

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 10



Prefijo global de Ruteo (Global Routing Prefix): Los 48 bits más significativos están designados como Global Routing Prefix el cual está asignado a un sistema autónomo específico. Los tres bits más significativos del Global Routing Prefix siempre serán 001.

Dirección de Enlace Local (Link-Local Address)

La dirección IPv6 auto-configurada es conocida como dirección de enlace local o Link-Local address. Esta dirección siempre empieza con FE80. Los primeros 16 bits de la dirección Link-Local siempre están establecidos como 1111 1110 1000 0000 (FE80).

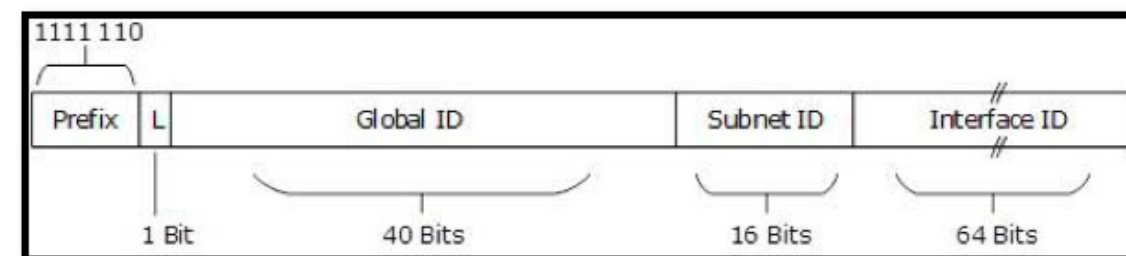
Las direcciones de enlace local son usadas para comunicación entre hosts IPv6 en un enlace (segmento de broadcast) solamente. Estas direcciones no son enrutables por tanto un router nunca enviará estas direcciones fuera del enlace.

Dirección Local Única (Unique-Local Address)

Este tipo de dirección IPv6 que, aunque está pensada globalmente como única, debe ser utilizada en comunicaciones locales. Esta dirección tiene la segunda mitad del ID de interfaz y la primera mitad se divide entre Prefijo, Bit local, Global ID y el ID subred.

Gráfico 9: Unique-Local Address

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 11



El prefijo siempre será establecido en 1111 1110. El bit L, si está establecido en 1 significa que la dirección fue asignada localmente. Hasta ahora, si L = 0, no tiene significado alguno. Por tanto, la dirección local única (Unique-Local Address) siempre empieza con "FD".

Ámbito de direcciones IPv6 de unicast:

El ámbito de las direcciones de enlace local está limitado al segmento. Las direcciones Locales Únicas están pensadas de manera global pero no son ruteadas hacia Internet, limitando su ámbito únicamente al perímetro organizacional. Las direcciones Globales Unicast son mundialmente únicas y reconocibles. Estas componen esencialmente el direccionamiento de Internet versión 2.

Gráfico 10: Ámbito de direcciones IPv6 de unicast

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 11



Direccionamiento Especial

La versión 6 tiene una estructura ligeramente más compleja de direcciones IP que la versión IPv4. IPv6 ha reservado algunas direcciones y notación de direcciones para propósitos especiales.

Como se muestra en la tabla anterior, la dirección 0:0:0:0:0:0:0:0/128 no especifica nada y se dice que es una dirección inespecífica. Después de simplificar todos los ceros se compacta a :: / 128.

En IPv4, la dirección 0.0.0.0 con máscara 0.0.0.0 representa la ruta por defecto. El mismo concepto es aplicado a IPv6, donde la dirección 0:0:0:0:0:0:0:0 con máscara de red todos ceros, representa la ruta por defecto. Después de aplicar la regla de simplificación de IPv6 esta dirección se comprime a ::/0.

La dirección de loopback en IPv4 está representada por la serie de 127.0.0.1 a 127.255.255.255. Pero en IPv6, solo la dirección 0:0:0:0:0:0:0:1/128 representa al loopback. Después de simplificar la

dirección de loopback, puede ser representada como ::1/128.

CABECERA IPv6

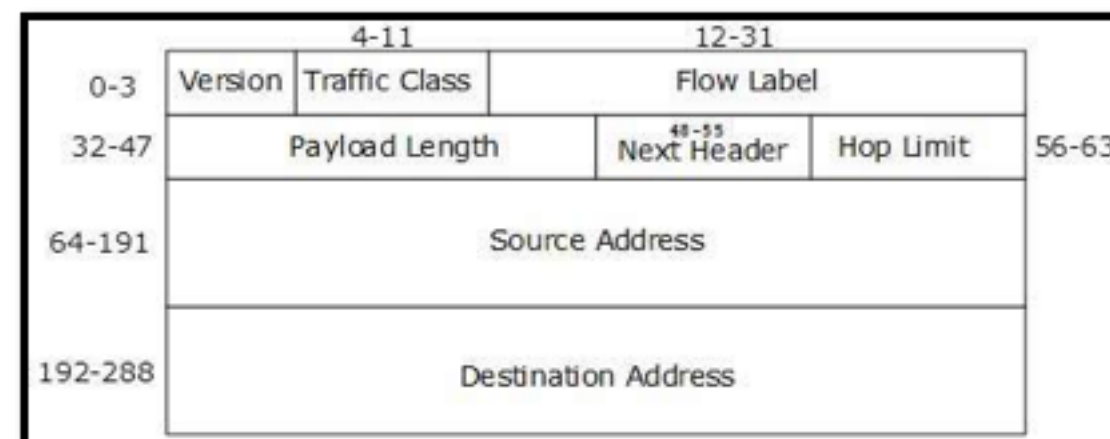
Lo mejor del nuevo protocolo IPv6 es sin duda alguna su cabecera. Las direcciones IPv6 son 4 veces más grandes que IPv4, pero la cabecera de IPv6 es tan solo dos veces mayor que la de IPv4. Las cabeceras de IPv6 tienen un segmento de cabecera fijo y cero o más cabeceras opcionales, las cuales son llamadas extensiones. Toda la información necesaria la cual es esencial para el router, se encuentra en la cabecera fija. Las cabeceras de extensión contienen información opcional la cual ayuda al router a entender cómo manejar el flujo de paquetes.

Cabecera fija

La cabecera fija de IPv6 tiene 40 bytes de longitud y contiene la siguiente información:

Gráfico 11: Cabecera fija de IPv6

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 14



Versión (Version): Tiene una extensión de 4 bits. Representa la versión del protocolo de Internet. Ejemplo 0110: Versión 6.

Clase de Tráfico (Traffic Class): Tiene una longitud de 8 bits. Estos 8 bits están divididos en dos partes. Los 6 bits más significativos son usados para definir el tipo de servicio, el cual dice al router qué servicios deben ser provistos a dicho paquete. Los 2 bits menos significativos son usados por ECN (Explicit Congestion Notification – Notificación Explícita de Congestión).

Etiqueta de Flujo (Flow Label): Longitud de 20 bits. Esta etiqueta es utilizada para mantener el flujo secuencial de los paquetes pertenecientes a una comunicación. La fuente etiqueta la secuencia la cual ayuda al router a identificar que el paquete pertenece a un determinado flujo de información. Este campo ayuda a evitar el re-ordenamiento de los paquetes de datos. Está diseñado para transmisión de paquetes RTP (Real Time Protocol – Protocolo de Tiempo Real) como VoIP, Telefonía IP, Video Conferencia, Video Vigilancia IP entre otros.

Longitud de la carga (Payload Length): Su extensión es de 16 bits. Este campo es utilizado para decir a los routers cuánta información contiene un determinado paquete en su carga. La carga está compuesta de cabeceras de extensión y datos de capa alta. Con 16 bits, hasta 65535 bytes pueden ser indicados, pero, si la cabecera de extensión contiene una cabecera Hop-by-Hop entonces la carga puede exceder esos 65535 bytes y su campo será establecido en cero.

Siguiente encabezado (Next Header): Tiene una longitud de 8

bits. Este campo es utilizado para indicar tanto el tipo de cabecera de extensión o, cuando la cabecera de extensión no está presente entonces indica la PDU (Packet Data Unit) de capa superior. Los valores de los tipos de PDU de capa superior son los mismos que en el protocolo IPv4.

Límite de saltos (Hop Limit): Este campo es utilizado para detener el bucle infinito de un paquete dentro de la red. Es lo mismo que TTL (Time To Live) en el protocolo IPv4. El valor de límite de saltos es disminuido cada vez que este paquete pasa a través de un enlace (router / salto). Cuando el campo llega a cero, el paquete es descartado.

Dirección Fuente (Source Address): Tiene una longitud de 128 bits. Este campo indica la dirección de quien origina el paquete.

Dirección Destino (Destination Address): Posee una extensión de 128 bits. Provee la dirección del destinatario del paquete.

Cabeceras de Extensión

En IPv6, la Cabecera Fija contiene únicamente información que es necesaria y evita información que no es requerida o poco utilizada. Toda esa información es colocada entre la Cabecera Fija y la Cabecera de Capa Superior en forma de Cabecera de Extensión. Cada Cabecera de Extensión es identificada a través de un valor distintivo.

Cuando las cabeceras de extensión son utilizadas, los campos Next Header de las cabeceras fijas IPv6, apuntan a la cabecera de extensión. Si existe una o más cabeceras de extensión, entonces el primer campo Next Header de la cabecera de extensión apuntará al segundo campo y así sucesivamente. El último campo Next Header

de la cabecera de extensión apunta a la cabecera de capa superior.
Así, todos los encabezados apuntarán al siguiente, en una forma lista
enlazada.

CAPÍTULO III

COMUNICACIÓN

En IPv4, un host el cual quiere comunicarse con algún otro host en la red, necesita primero tener una dirección IP adquirida ya sea a través de un servidor DHCP o configurada de manera manual. Tan pronto como el equipo disponga de la dirección IP válida, será capaz de hablar con cualquier host en la red. Para comunicarse a través de capa 3, un host también debe conocer la dirección IP del otro host. La comunicación sobre un enlace, es establecida por medio de la dirección MAC embebida en el hardware. Para conocer la dirección MAC de los equipos de los cuales se sabe la dirección IP, un host envía un broadcast de ARP (Address Resolution Protocol) y como respuesta, el host de interés envía su dirección MAC.

En IPv6, no existe un mecanismo de broadcast. No es obligatorio para un host habilitado para IPv6 el obtener una dirección IP de un servidor DHCP o de manera manual, ya que este puede auto-configurar su propia dirección IP. Entonces, a la pregunta de ¿cómo un host puede comunicarse con otros en una red habilitada para IPv6?, podemos responder que el protocolo ARP ha sido reemplazado por ICMPv6 Neighbor Discovery Protocol (Protocolo de Descubrimiento de Vecinos).

Protocolo de Descubrimiento de Vecinos (Neighbor Discovery Protocol)

Un host en una red habilitada para IPv6 es capaz de auto-configurarse a sí mismo con una dirección única de enlace local. Tan pronto como tiene una dirección IPv6, este se une a numerosos grupos multicast. Todas las comunicaciones relacionadas a ese segmento suceden en esas direcciones multicast solamente. Un host atraviesa algunos estados en IPv6, a saber:

- **Solicitud de Vecino (Neighbor Solicitation):** Después de configurar todas las IPv6 sea de manera manual, con un servidor DHCP o, por auto-configuración, el host envía un mensaje de “solicitud de vecino” a la dirección multicast FF02::1/16 para todas sus direcciones IPv6. Con esto determina que nadie ocupa una misma dirección.
- **DAD (Duplicate Address Detection – Detección de Dirección Duplicada):** Cuando un host no escucha nada en el segmento respecto de su mensaje de solicitud de vecino, este asume que no existen direcciones duplicadas en el segmento.
- **Anuncio de Vecino (Neighbor Advertisement):** Después de asignar la dirección a sus interfaces, levantarlas y ponerlas en servicio, el host una vez más envía un mensaje de Neighbor Advertisement diciéndoles a todo el resto de hosts, que él se ha asignado esas IPv6 a sus interfaces.

Una vez que el host ha terminado la configuración de sus direcciones IPv6, este realiza las siguientes tareas:

- **Solicitud de Router (Router Solicitation):** Un host envía un paquete multicast (FF02::2/16) de solicitud de router en su segmento para conocer si existe algún equipo router. Esto ayuda al host a

configurar la dirección del posible router de segmento como su dirección de puerta de enlace predeterminado. Si este router deja de trabajar, el host puede cambiarse a otro router y convertirlo en su nueva puerta de enlace predeterminado.

- Anuncio de Router (Router Advertisement): Cuando un router recibe un mensaje de solicitud de router, este responde de regreso, para anunciarse ante la solicitud del host, en ese enlace.
- Redirección (Redirect): Esta puede ser la situación en que un router recibe una Solicitud de Router, pero él sabe que no es la mejor puerta de enlace para ese host. En esta situación, el router avisa a ese host que existe disponible un mejor router de “next-hop (próximo salto)”.

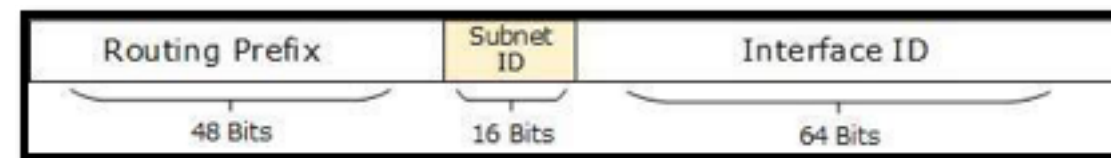
2.10 SUBNETTING

En IPv4, las direcciones fueron creadas en clases. Las direcciones IPv4 con clase claramente definen el bit usado para los prefijos de red y el bit usado por los hosts en la red. Para hacer subredes en IPv4 nos valemos de la máscara de clase por defecto la cual permite pedir prestado bit de hosts para ser utilizados como bits de subred. Esto resulta en múltiples subredes, pero menor número de hosts por subred.

Las direcciones IPv6 utilizan 128 bits para representar una dirección la cual incluye bits que serán usados para hacer subredes. La segunda mitad de la dirección (los 64 bits menos significativos) es siempre utilizada para Host únicamente.

Gráfico 13: Subredes de IPv6

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 19



Los 16 bits del campo Subnet (Subred) equivalen a una red clase B de IPv4. Utilizando estos bits de subred, una organización puede tener más de 65 mil subredes las cuales son más que suficientes.

Por lo tanto, el prefijo de enrutamiento /64 y la porción de host es de 64 bits. Sin embargo, se puede dividir en subredes más allá de los 16 bits del identificador de subred (Subnet ID), tomando prestado bits de los hosts, pero, se recomienda que los 64 bits siempre sean utilizados para las direcciones de hosts porque la auto-configuración requiere de 64 bits.

Las subredes en IPv6 trabajan bajo el mismo concepto de enmascaramiento de subredes de longitud variable en IPv4.

El prefijo /48 puede ser usado en un organización proveyendo los beneficios de tener hasta subredes de prefijo /64, las cuales serían 65535 subredes, cada una teniendo 264 host. Un prefijo /64 puede ser asignado a una conexión punto – punto, donde hay solo dos hosts en un enlace.

CAPÍTULO IV

IPv4 A IPv6

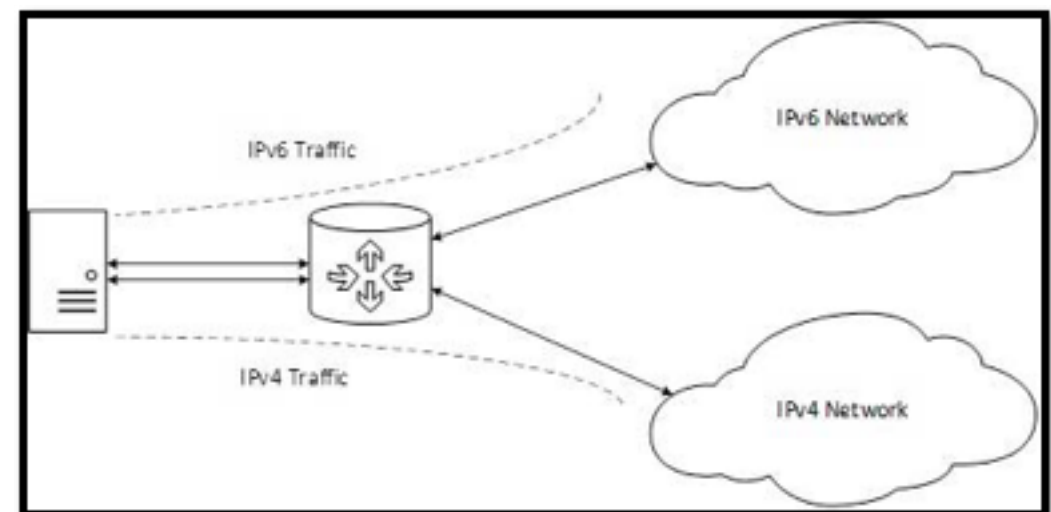
Un problema en la transición completa de IPv4 a IPv6 es que IPv6 no es compatible hacia atrás. A diferencia de la implantación de otras nuevas tecnologías, donde lo nuevo es compatible hacia atrás, con los sistemas antiguos para que se pueda seguir trabajando sin cambios mayores, en IPv6 esto no ocurre; lo que resulta en una situación donde un sitio es IPv6 o no lo es.

Doble Pila de Router (Dual Stack Router)

Un router puede ser instalado con direcciones IPv4 e IPv6 configuradas en sus interfaces y apuntando a la red de un esquema IP relevante.

Gráfico 14: Doble pila de router

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 20



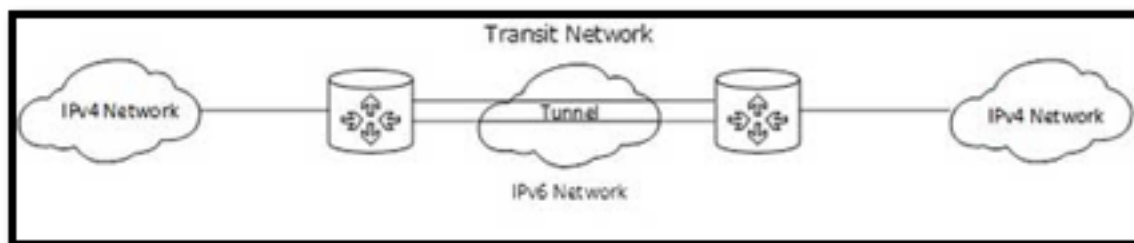
En el gráfico anterior, un servidor el cual tiene tanto direcciones IPv4 cuanto IPv6 configuradas podrá comunicarse con los hosts de las redes IPv4 e IPv6 con la ayuda de Dual Stack Router. El Dual Stack Router puede comunicarse con ambas redes y proveer un medio para que los hosts accedan al servidor sin cambiar sus respectivas versiones de IP.

Tunelizado

En un escenario donde diferentes versiones de IP existen en una ruta intermedia o red de tránsito, el tunelizado provee la mejor solución donde los datos de usuario pueden pasar a través de sin importar la versión de IP.

Gráfico 15: Tunelizado

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 21.



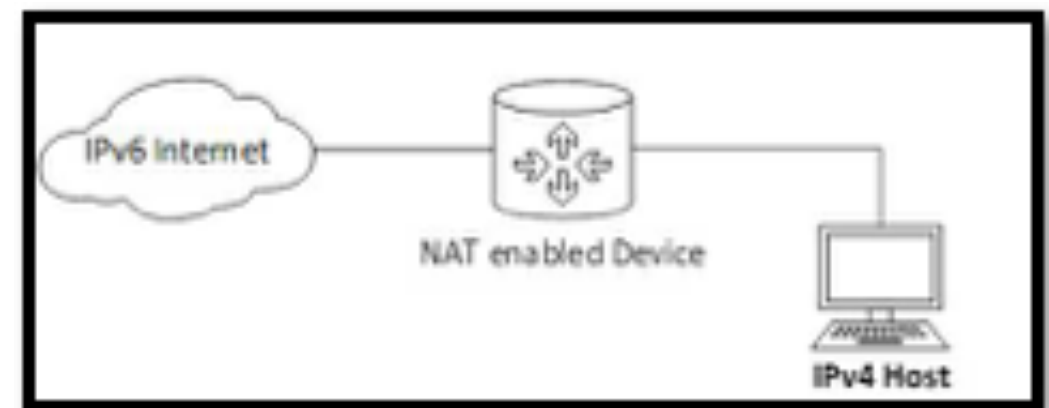
El diagrama anterior demuestra como dos redes IPv4 remotas pueden comunicarse a través del túnel, donde la red de tránsito es un IPv6.

Protocolo de Traducción NAT

Otro método importante en la transición hacia IPv6 es por medio de NAT-PT (Network Address Translation – Protocol Translation) habilitado en el dispositivo. Con la ayuda del dispositivo NAT-PT, la conversión sucede entre los paquetes IPv6 e IPv4 y viceversa. Mirar el diagrama siguiente:

Gráfico 16: NAT-PT

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 21



Un host con una dirección IPv4 envía una petición a un servidor habilitado para IPv6 en Internet el cual no entiende direcciones IPv4. En este escenario un dispositivo NAT-PT puede ayudar a comunicarlos. Cuando el host IPv4 envía una solicitud de paquetes al servidor IPv6, el dispositivo/router NAT-PT, modifica el paquete IPv4, remueve la cabecera de IPv4 y adiciona una cabecera IPv6 y pasa a través de Internet. Cuando el servidor con IPv6 responde hacia el host IPv4, el router hace el trabajo inverso.

MOVILIDAD

Cuando un host es conectado en una red, este adquiere una dirección IP y toda la comunicación sucede utilizando esa dirección IP en el enlace. Tan pronto como ese mismo host cambia físicamente de ubicación, esto es, se mueve a un área diferente (subred, red, enlace), su dirección IP cambia y toda la comunicación que sucedía en el host usando la dirección IP antigua, se cae.

La movilidad de IPv6 provee un mecanismo el cual equipa a un host con la habilidad de recorrer entre diferentes enlaces sin perder ninguna conexión ni tampoco su dirección IP.

Múltiples entidades están envueltas en esta tecnología:

- **Nodo Móvil (Mobile Node):** El dispositivo que requiere de movilidad IPv6.
- **Home Link:** Este enlace es configurado con el prefijo de subred de la red de inicio (home) y esta es donde el dispositivo móvil IPv6 adquiere su dirección inicial (home address).
- **Home Address (Dirección Local Permanente):** Es la dirección que el nodo móvil adquiere de Home Link. Si el dispositivo móvil permanece en el mismo Home Link, la comunicación entre varias entidades ocurre como usualmente se efectúa.
- **Home Agent:** Este es un router el cual actúa como registrador de los Nodos Móviles. El Home Agent está conectado al Home Link y mantiene información sobre todos los Nodos Móviles, sus Home Address y sus direcciones IP actuales.
- **Foreign Link (Enlace Foráneo):** Visto desde el lado del Nodo Móvil, cualquier otro enlace que no se a su Home Link.

- **Care-of Address (Dirección Dinámica de Visita):** Cuando un nodo móvil se adjunta a un enlace foráneo (foreign link), este adquiere una nueva dirección IP de esa subred del enlace externo (foráneo). El router Home Agent mantiene la información de la dirección original (Home Address) y de la dirección dinámica (Care-of Address). Múltiples direcciones dinámicas pueden ser asignadas a un nodo móvil, pero en cualquier instancia, solo una dirección dinámica (Care-of Address) está ligada a la dirección local (Home Address).

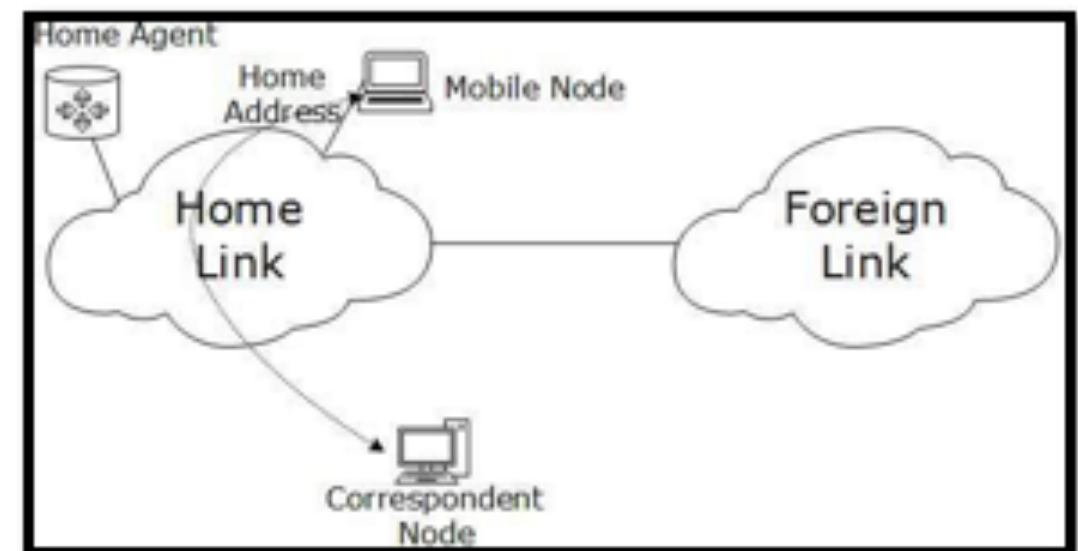
- **Nodo Correspondiente:** Cualquier dispositivo IPV6 el cual intente establecer comunicación con el Nodo Móvil.

2.12.1 Operación de Movilidad

Cuando un Nodo Móvil permanece en su enlace local (Home Link), todas las comunicaciones suceden en su dirección local (Home Address). Como se muestra a continuación:

Gráfico 17: Nodo Móvil conectado al Home Link

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 23

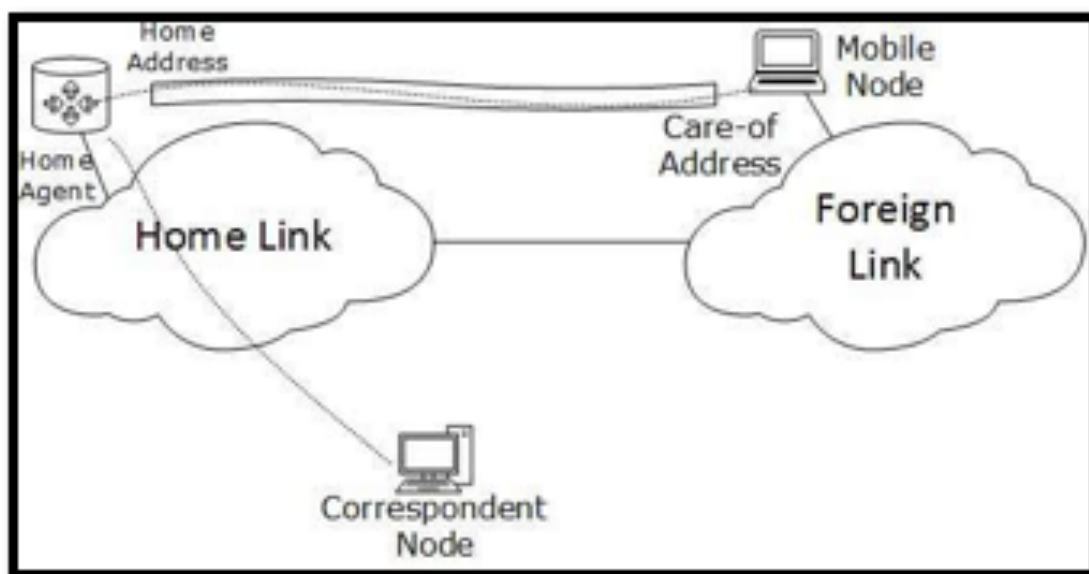


Cuando un Nodo Móvil deja su Home Link y se conecta a algún enlace externo (Foreign Link), la característica de movilidad de IPv6 entra en juego. Después de conectarse al enlace externo, el nodo móvil adquiere una dirección IPv6 desde el enlace externo. Esta dirección es llamada dirección dinámica (Care-of Address). El nodo móvil envía una petición de unión a su router Home Agent con la nueva dirección Care of Address. El Home Agent une la dirección local (Home Address) del nodo móvil con la dirección dinámica (Care of Address), estableciendo un túnel entre ambas.

Cuando un nodo correspondiente trata de establecer conexión con el nodo móvil (en su Home Address), el Home Agent intercepta el paquete y lo reenvía a la dirección dinámica del nodo móvil sobre el túnel que fue previamente establecido.

Gráfico 18: Nodo Móvil conectado a un enlace externo

Recuperado de: Tutorials Point, Learn IPv6 Internet Protocol Version 6, Página 23



Optimización de Ruta

Cuando un Nodo Correspondiente inicia comunicación enviando paquetes al nodo móvil en su Home Address, estos paquetes son tunelizados hacia el nodo móvil por el Home Agent. En el modo de optimización de ruta, cuando el nodo móvil recibe un paquete desde el No Correspondiente, este no envía una respuesta al Home Agent. En lugar de eso, envía directamente su paquete al nodo correspondiente usando su Home Address como dirección fuente. Este modo es opcional y no es usado por defecto.

RUTEO

Los conceptos de ruteo permanecen iguales en el caso de IPv6, pero, casi todos los protocolos de ruteo han sido redefinidos. Se ha mencionado en el apartado Protocolo de Comunicación IPv6, como los Hosts hablan a las puertas de enlace predeterminado (gateways). El ruteo es el proceso de enviar datos ruteables (ya que existen datos que no se pueden rutear como se había mencionado antes) escogiendo el mejor camino entre muchos disponibles. Un router es un dispositivo el cual reenvía datos los cuales no están explícitamente destinados hacia él.

Existen dos formas de protocolos de enrutamiento, a saber:

Protocolo de Enrutamiento del Vector Distancia: Un router que corre el protocolo de vector distancia anuncia sus rutas conectadas y aprende nuevas rutas de sus vecinos. El costo de enrutamiento para alcanzar un destino es calculado por medio de los saltos entre la fuente y el destino. Un router generalmente se basa en sus vecinos para la selección de la mejor dirección, también conocido como “enrutamiento por rumores”. RIP y BGP son protocolos de vector

distancia.

Protocolo de Enrutamiento de Estado del Enlace: Este protocolo reconoce el estado de un enlace y lo anuncia a sus vecinos. La Información sobre los nuevos enlaces se aprende de los routers vecinos. Después de que toda la información de enrutamiento se ha convergido, el Protocolo de Enrutamiento de Estado del Enlace utiliza su propio algoritmo para calcular la mejor ruta para todos los enlaces disponibles. OSPF e IS-IS son los protocolos de enrutamiento de estado de enlace y ambos utilizan el algoritmo de Dijkstra para la ruta más corta primero.

Los protocolos de enrutamiento pueden ser divididos en dos categorías:

Protocolos de Enrutamiento Interno: Los protocolos en esta categoría son usados en un sistema autónomo o en una organización para distribuir las rutas entre todos los routers dentro del perímetro. Ejemplos: RIP, OSPF.

Protocolo de Enrutamiento Exterior: Un protocolo de enrutamiento exterior distribuye información de enrutamiento entre dos sistemas autónomos diferentes u organizaciones. Ejemplos: BGP.

Protocolos de Enrutamiento

Existen diversos protocolos de enrutamiento. Se tratarán a continuación los más utilizados en las organizaciones:

RIPng: Protocolo de Información de Enrutamiento de próxima generación (Routing Information Protocol Next Generation). Se trata de un protocolo de enrutamiento interior y es corresponde a los del tipo vector de distancia. RIPng se ha actualizado para soportar IPv6.

OSPFv3: La versión 3 de OSPF es un protocolo de enrutamiento interior el cual ha sido modificado para soportar IPv6. Este es un

protocolo de Estado – Enlace y usa el algoritmo de Dijkstra para calcular la mejor ruta.

BGPv4: Sus siglas representan Border Gateway Protocol. Es la única norma de Protocolo Gateway Exterior de código abierto disponible. BGP es un protocolo de vector de distancia que tiene como métrica de cálculo un Sistema Autónomo, en lugar de un número de routers como salto. BGPv4 es una actualización de BGP para soportar el enrutamiento de IPv6.

Protocolos modificados para soportar IPv6

ICMPv6: El protocolo de control de mensajes de Internet, en su versión 6, es una actualización de ICMP para soportar los requerimientos de IPv6. Este protocolo es utilizado para funciones de diagnóstico, mensajes de información y errores y, propósitos estadísticos. El protocolo de Descubrimiento de Vecinos ICMPv6 reemplaza a ARP y ayuda a descubrir vecinos y routers en el enlace.

DHCPv6: Es una implementación de DHCP. Aunque los hosts habilitados para IPv6 no requieren ningún servidor DHCPv6 para adquirir una dirección IP ya que se pueden auto-configurar. Tampoco necesitan del DHCPv6 para localizar el servidor DNS porque el DNS puede ser descubierto y configurado a través de ICMPv6. Todavía el servidor DHCPv6 puede ser utilizado para proveer de toda esa información.

DNS: No existe una nueva versión de DNS pero ahora está equipada con extensiones para proveer soporte para consultas de direcciones IPv6. Un nuevo registro AAAA (cuádruple A) ha sido añadido para responder a los mensajes consulta de IPv6. Ahora el DNS puede responder tanto con versiones IP 4 y 6 sin ningún cambio

en el formato de consulta.

FUTURO DE IPv6

IPv6 habilita a la versión 2 de Internet para remplazar a la Internet IPv4 de nuestros días. Cuando Internet fue lanzado con IPv4, los países desarrollados como Estados Unidos y algunos de Europa tomaron grandes espacios de IPv4 para desarrollar Internet en sus respectivos países teniendo presente las necesidades futuras. Pero Internet explotó en todas partes alcanzando y conectando a todos los países del mundo, aumentando así el requerimiento de espacio de direcciones IPv4.

Como resultado, hasta el día de hoy los Estados Unidos y Europa tienen mucho del espacio de direcciones IPv4 y países como India y China están obligados a hacer frente a su necesidad de espacio IP por medio del despliegue de IPv6.

La mayor parte del despliegue de IPv6 se está haciendo fuera de los Estados Unidos y Europa. India y China se están adelantando para cambiar la totalidad de su espacio a IPv6. China ha anunciado un plan de despliegue de cinco años llamado: China, Internet de próxima generación.

Después del 6 de junio de 2012 todos los principales ISP se cambiaron a IPv6 y el resto de ellos todavía se está moviendo.

IPv6 proporciona un amplio espacio de direcciones y está diseñado para ampliar los servicios de Internet de hoy en día. La versión 2 de Internet, rica en características compatibles con IPv6, puede entregar más de lo esperado.

IMPLEMENTACIÓN DE CALIDAD DE SERVICIO EN IPv6

De manera conceptual, se define a la Calidad de Servicio (QoS) como la capacidad de una red para proporcionar diversos niveles de servicio a los diferentes tipos de tráfico. A través de los mecanismos de QoS se asegura la correcta entrega de los paquetes IP, dando prioridad a las aplicaciones de Misión Crítica, donde se comparten a la par los recursos con aplicaciones no críticas. QoS hace la diferencia al proveer un uso eficiente de los recursos en caso de presentarse congestión en la red, seleccionando un tráfico específico de ésta, priorizándolo según su importancia relativa, y utilizando métodos de control y evasión de la congestión para darles un tratamiento preferencial. Implementando QoS en una red, se logra un rendimiento de ésta más predecible y una utilización de ancho de banda más eficiente. (Palacias, 2011).

El Protocolo de Internet en su mayor parte, trata a todos los paquetes por igual, los cuales se remiten con el tratamiento de mejor esfuerzo y ninguna garantía para la entrega a través de la red. TCP (Transmission Control Protocol) añade las confirmaciones de entrega, pero no tiene opciones para controlar parámetros como el retraso o la asignación de ancho de banda. QoS ofrece una mayor cantidad de opciones de red, basadas en políticas, para dar prioridad a la entrega de la información. Implementaciones de IPv4 e IPv6 existentes utilizan similares capacidades de calidad de servicio, tales como servicios diferenciados y servicios integrados, para identificar y priorizar las

comunicaciones basadas en IP durante los períodos de congestión de la red. Dentro de la cabecera IPv6 dos campos pueden ser utilizados para QoS, la clase de tráfico y la etiqueta de flujo. El nuevo campo de etiqueta de flujo y el campo ampliado Clase de tráfico en la cabecera principal IPv6 permiten la diferenciación más eficiente y detallada de diversos tipos de tráfico. El nuevo campo Etiqueta de flujo puede contener una etiqueta de identificación o dar prioridad a un cierto flujo de paquetes tal como voz sobre IP (VoIP) o videoconferencia, los cuales son sensibles a la entrega oportuna. IPv6 QoS es todavía un trabajo en progreso y la seguridad se deberá incrementar considerablemente en esta etapa de desarrollo.

La capa de Red del protocolo TCP/IP, tanto de IPv4 cuanto de IPv6, fue intencionalmente diseñada sin ninguna de las características que normalmente se asocian con la calidad de servicio, tales como los controles de admisión, garantías de recursos y entrega sin pérdidas. QoS en redes TCP / IP o sobre la Internet es un concepto algo impreciso, el cual puede tener diferentes significados en función de la óptica aplicada. Puede significar, por ejemplo:

Proporcionar a un usuario ciertos niveles de disponibilidad general, de rendimiento, baja latencia, valor máximo de pérdida de paquetes, o incluso cierto nivel de seguridad.

El tratamiento de diferentes tipos de tráfico de manera desigual, de acuerdo con el contenido: audio o vídeo en tiempo real, por ejemplo, requiere un alto rendimiento y una constante y reducida latencia. Este tipo de tráfico puede tolerar mínimas pérdidas, mientras que la transferencia de archivos puede tolerar demora, pero no pérdidas de ningún tipo.

Muchos aspectos de la ingeniería de calidad de servicio dependen de las tecnologías que se ejecutan en múltiples capas, tales como MPLS (MultiProtocol Label Switching) y ATM (Asynchronous Transfer Mode). El trabajo de la IETF en QoS para TCP / IP se inició con Servicios Integrados (IntServ), que fue diseñado para proporcionar garantías de calidad de servicio. Intserv ha sido reemplazado por los Servicios Diferenciados (DiffServ), el cual simplemente reconoce que diferentes tipos de tráfico tienen diferentes requisitos de calidad de servicio y necesitan ser marcados de manera adecuada. El protocolo de señalización, es decir, el protocolo utilizado para la especificación y configuración de calidad de servicio, para establecer las peticiones de QoS es el Protocolo de Reserva de Recursos RSVP (Resource Reservation Protocol). El direccionamiento IPv6 de extremo a extremo permite la posibilidad de implementar servicios que son complicados de hacerlo con NAT, así como también el uso de extremo a extremo de DiffServ y RSVP.

Muchos de estos servicios poseen contenido multimedia y aplicaciones en tiempo real, por lo que es probable que se convierta en un tema más importante con el uso generalizado de la QoS en IPv6.

La noción de la mejora de QoS siempre estará ligada a IPv6. De hecho, IPv6 fue diseñado para soportar ciertas mejoras de calidad de servicio, pero no todas ellas han sido completamente especificadas o aplicadas.

Clasificación

El concepto de clasificación es utilizado para separar paquetes de datos, basado en determinadas características como pueden ser la dirección de origen y destino. Esto se logra predefiniendo patrones en

el campo ToS de 8 bits del encabezado IP. La clasificación es uno de los puntos más importantes del concepto Calidad de Servicio ya que sin esta característica aplicada a la red, todos los paquetes IP serían tratados de la misma manera.

La clasificación también puede basarse en información de protocolos de nivel superior y en otros descriptores de tráfico tales como:

- Interfaz de ingreso

- Valor CoS en la trama 802.1p

- Valor DSCP de la cabecera del paquete IP

- Valor MPLS EXP de la cabecera MPLS

De manera general, el proceso de clasificación consiste en identificar cada tipo de tráfico en la red y categorizarlo dentro de clases. De hecho, uno de los mejores métodos y de uso cada vez más general es por línea de comandos o MQC (Modular QoS Command-Line). Este método puede clasificar el tráfico usando además de listas de control de acceso (ACL), descriptores de tráfico, tipo de aplicación, al utilizar NBAR (Network Based Application Recognition).

Aplicaciones de Misión Crítica incluido ERP y aplicaciones workforce optimization pueden ser identificadas inteligentemente y clasificadas usando NBAR. Una vez que las aplicaciones de misión crítica son clasificadas estas pueden garantizarse una mínima cantidad de ancho de banda, política de ruteo, y marcadas para tratamiento preferencial. Aplicaciones que no son críticas incluyendo aplicaciones de juegos sobre Internet y el compartir archivos MP3 también pueden ser clasificadas usando NBAR y marcadas para servicio de mejor esfuerzo, crearles una política o bloquearlas según se requiera.

Marcaje

El campo ToS del encabezado del paquete puede ser reemplazado por los enrutadores con un valor relevante a las políticas de QoS definidas en la red. Esta acción sobre un paquete se denomina marcado. (Shenker. 1994).

El campo Etiqueta de Flujo de 20 bits en la cabecera IPv6 se agrega para permitir el etiquetado de paquetes que pertenecen a flujos de tráfico particulares y puede ser usado por el origen para etiquetar secuencias de paquetes para las cuales solicita un manejo especial por parte de los enrutadores IPv6, tal como la calidad de servicio no estándar o el servicio en tiempo real.

El campo de ocho bits Clase de Trabajo en la cabecera IPv6 es utilizado por los nodos origen y/o enrutadores intermedios para identificar y distinguir entre las diferentes clases o prioridades de paquetes IPv6. Su función es similar al campo ToS en IPv4 (Palacios. 2011).

Luego del proceso inicial de clasificación, el paquete es marcado de tal manera que se lo puede identificar dentro de una clase. A posteriori, el paquete será identificado por su marca y se le otorgará un nivel de servicio acorde.

Asignación de Políticas

En la actualidad existen diversas técnicas que logran estandarizar y cumplir con los requerimientos de QoS de las aplicaciones RTP (Telefonía IP, e-commerce, videoconferencia), tales como IntServ (Servicios Integrados), DiffServ (Servicios Diferenciados) y MPLS (Multiprotocolo de conmutación de etiquetas), todas estas con sus ventajas y desventajas.

Dentro de las técnicas y mecanismos más difundidos para la

asignación de políticas de calidad de servicio se encuentran tres, a saber:

Best Effort (Mejor Esfuerzo): Esta técnica define que la red en su conjunto hará todo lo posible por entregar adecuadamente los paquetes en el destino, sin embargo, no existe garantía alguna de que esto ocurra. Este modelo se lo utiliza en aplicaciones como HTTP y FTP.

Servicios Integrados (IntServ): Integrated Services brinda a las aplicaciones un nivel garantizado de servicio, negociando parámetros de red end-to-end. La aplicación per se solicita el nivel de servicio acorde a sus necesidades con el propósito de operar apropiadamente y se basa en la QoS para que se reserven los recursos de red necesarios antes de que la aplicación empiece a operar. La arquitectura de Servicios Integrados o IntServ corresponde a la RFC 1633. El modelo se sustenta en los siguientes supuestos:

Con el propósito de satisfacer los requerimientos de las aplicaciones, los recursos se deben gestionar de forma directa y explícita. Esto implica el uso de mecanismos para control de admisión y reservación de recursos.

Internet debe ser la infraestructura común para el tráfico normal y de tiempo real. Montar una nueva red para el tráfico RTP resultaría demasiado complejo. Por tanto, se debe unificar la pila de protocolos para cualquier tipo de tráfico, es decir, IP debe ser utilizado también para el transporte de datos de tiempo real. (Palacios. 2011).

Cada flujo se debe atender independientemente y no puede influenciar a otros. Adicional a la técnica de Best Effort, la arquitectura define dos clases más de servicios, estos son: Servicio Garantizado y Servicio de

Carga Controlada, mismos que definen el tratamiento que recibirán los flujos a lo largo del camino. Integrated Services necesita que los recursos imperiosos para satisfacer los requerimientos de una aplicación o servicio se reserven sobre el trayecto con anticipación, para lo cual es necesario el Protocolo de Reservación de Recursos (RSVP). Éste usa un conjunto de mensajes de señalización para transportar información sobre los requerimientos y propiedades de cada flujo, la cual se utiliza para mantener tablas de estado en cada uno de los nodos, generando así un alto tráfico de señalización y ocupación de los recursos en los dispositivos. (Flannagan. 2001).

Servicios Diferenciados (DiffServ): Este método abarca diversas herramientas de clasificación y mecanismos de encolamiento, los cuales proveen a ciertas aplicaciones de prioridades determinadas sobre el resto de paquetes en la red. La teoría de los servicios diferenciados se la puede hallar descrita en el RFC 2474.

El principal problema que se presenta al tratar de llevar a la práctica el modelo de servicios integrados en Internet es la potencial existencia de un número infinito de tipos de tráfico diferentes en la red. Como consecuencia de ello, cada dispositivo de encaminamiento debe almacenar, en relación a cada flujo particular, la información de estado necesaria para proporcionarle la calidad de servicio que le ha sido garantizada. Mediante el modelo de servicios diferenciados se pretende acomodar las diversas expectativas de los usuarios y los requisitos heterogéneos de las aplicaciones, pero dentro de unos límites, de modo que se evite una sobrecarga por almacenamiento de la información de estado y por el procesamiento exigido, (España, 2003). Dentro del protocolo IPv6, los octetos de la cabecera clase de tráfico

son reinterpretados, pasando a denominarse campo de servicios diferenciados. Los paquetes IP son clasificados a la entrada en la red entre un conjunto acotado de clases de servicios, identificándose a qué clase pertenece un paquete mediante el código denominado código de servicios diferenciados, que se inserta en el campo de servicios diferenciados de su cabecera. Este proceso como se había descrito en apartados anteriores, es denominado marcado. Una vez que los dispositivos de ruteo reciben un paquete, estos examinan dicho campo y le facilitan un tratamiento distinto según su clase.

El modelo de servicios diferenciados es, de algún modo, un esquema de prioridades, en el cual paquetes pertenecientes a clases diferentes son tratados con distinta prioridad. Comparándolo con el protocolo IPv4 y su reducido espacio de opciones del modelo tipo de servicio, el cual no es capaz de hacer frente al crecimiento en número y variabilidad de los servicios previstos en la futura Internet, se debe dejar por sentado que la arquitectura de servicios diferenciados de IPv6 permite superar esta limitación.

Manejo de la congestión en la red

Manejar o administrar la congestión es un término que engloba diversas estrategias de encolamiento, con el propósito de tratar situaciones donde la demanda de recursos de las aplicaciones, excede el ancho de banda total que puede brindar la red. Entre los mecanismos de encolamiento encontramos los siguientes:

FIFO: Consiste en el tipo más simple de encolamiento. En este modelo, un búfer sencillo, retiene los paquetes salientes hasta que la interfaz de transmisión pueda enviarlos. Los paquetes se envían fuera de la interfaz en el mismo orden en que fueron llegando al búfer,

acordes al concepto FIFO (First Input First Output – Primero que Entra, Primero que Sale).

Cola de Prioridad (PQ): Es un modelo sencillo que ofrece tratamiento preferencial a los paquetes identificados. El modelo consta de cuatro colas: baja, normal, media y alta prioridad. Los paquetes que llegan a la interfaz se clasifican de acuerdo a dicha connotación. La salida de estas cuatro colas alimenta el búfer de transmisión de la interfaz. Los paquetes siempre se sirven desde las primeras colas de alta prioridad. El trabajo de este mecanismo resulta interesante en ambientes donde existe un tráfico importante, pero, no está optimizado para tratar las colas de menor prioridad, cayendo en la total falta de atención a estas últimas.

Cola Personalizada (CQ): Este tipo de encolamiento permite priorizar el tráfico sin los efectos laterales de agotamiento de las colas de baja prioridad, especificando el número de paquetes o bytes que deben ser atendidos para cada cola. Es posible crear hasta 16 colas para dar una categoría al tráfico, donde cada cola es atendida al estilo Round Robin (uno de los algoritmos de planificación de procesos más simples dentro de un sistema operativo que asigna a cada proceso una porción de tiempo equitativa y ordenada). El método de Cola Personalizada es usado para proporcionar a tráficos particulares de un ancho de banda garantizado en un punto de posible congestión, asegurando una porción fija del ancho de banda y permitiendo al resto del tráfico utilizar los recursos disponibles.

Weighted Fair Queuing (WFQ): El mecanismo de encolamiento WFQ asigna un peso determinado a cada flujo de forma que se genera un orden de tránsito en la cola de paquetes. Mediante los campos

discriminadores en TCP / IP (número de socket, tipo de protocolo, dirección de origen y destino) y por el ToS en el protocolo IP, se efectúa la mencionada ponderación. Weighted Fair Queuing crea una cola separada para cada tipo de tráfico y utiliza un valor predeterminado para la profundidad de la cola.

Modified Deficit Round Robin (MDDR): El método MDDR atiende a las colas que no están vacías, una tras otra en forma round robin. Cuando una cola es liberada, una cierta cantidad fija de datos se desencola y entonces el algoritmo atiende a la siguiente cola. MDDR hace el seguimiento del número de bytes de datos que fueron desencolados por encima del valor configurado, una vez que la cola ha sido atendida. En el siguiente paso, cuando la cola es atendida nuevamente, menos datos son desencolados para compensar el excedente de datos atendidos en el turno anterior. Como resultado, la cantidad promedio de datos atendidos, por cola, será muy cercano al valor configurado. Adicionalmente, MDDR mantiene una cola prioritaria que se atiende de manera preferencial, (Palacio, Salcedo & Lopez, 2011).

LLQ (Low Latency Queuing): Es un algoritmo que presenta aspectos importantes en el manejo de tráfico, ya que le da mayor prioridad a las aplicaciones sensibles al retardo, sin olvidar la cola de baja prioridad. Conmuta entre las diferentes colas, asignando más ancho de banda a las colas con mayor prioridad, logrando que estas lleguen en tiempo ideal a su destino final. La característica de LLQ brinda la posibilidad de especificar el comportamiento de baja latencia a las clases de tráfico tanto de misión crítica cuanto a las que no son de misión crítica. LLQ permite a los datos sensibles al retardo como

lo es el tráfico de voz, ser desencolado y enviarse primero (antes que los paquetes en las otras colas sean desencolados), brindando tratamiento preferencial a los datos sensibles al retardo sobre el resto de tráfico. Low Latency Queuing también introduce el concepto de sintonización del límite del anillo de transmisión. Antes de la introducción de LLQ, la máxima profundidad del anillo de transmisión no era un parámetro configurable por el usuario. De esta manera, algunas partículas podían acumularse en el anillo de transmisión sin límite, lo cual podía desembocar en una altísima e incontrolable latencia. La característica de LLQ permite a los usuarios limitar el número de partículas que pueden existir en el anillo de transmisión, bajando efectivamente la latencia incurrida por los paquetes situados en el anillo de transmisión.

Class-Based WFQ. El WFQ tradicional presenta algunas limitaciones de escalamiento, ya que la implementación del algoritmo se ve afectada a medida que el tráfico que corre por el enlace aumenta y llega a colapsar debido a la numerosa cantidad de flujos que se debe analizar. CBWFQ fue desarrollada para evitar estas limitaciones, tomando el algoritmo de WFQ y expandiéndolo, permitiendo la creación de clases definidas por el usuario, que permiten un mayor control sobre las colas de tráfico y asignación del ancho de banda. Algunas veces es necesario garantizar una determinada tasa de transmisión para cierto tipo de tráfico, lo cual no es posible mediante WFQ, pero sí con CBWFQ. Las clases que son posibles implementar con CBWFQ pueden ser determinadas según el protocolo ACL, valor DSCP, o interfaz de ingreso. Cada clase posee una cola separada, y todos los paquetes que cumplen el criterio definido para una clase

en particular son asignados a dicha cola. Una vez que se establecen los criterios para las clases, es posible determinar cómo los paquetes pertenecientes a dicha clase serán manejados. Si una clase no utiliza su porción de ancho de banda, otras pueden hacerlo. Se pueden configurar específicamente el ancho de banda y límite de paquetes máximos (o profundidad de cola) para cada clase. El peso asignado a la cola de la clase es determinado mediante el ancho de banda asignado a dicha clase, (Álvarez & Gonzalez, 2005).

Técnicas para evitar la congestión

Es procedente mencionar que los métodos de encolamiento descritos con anterioridad, no solucionan el problema de congestión en la red, sino que establecen determinadas reglas para que el tráfico más sensible tenga cierta prioridad sobre el resto de tráfico dentro de la red. Las técnicas para evitar congestión, por su parte, monitorean el flujo de tráfico de la red con el fin de anticipar y minimizar su efecto.

Random Early Detection (RED): Esta técnica monitorea el tamaño de la cola y, cuando ésta llega a un nivel previamente especificado, aleatoriamente escoge flujos TCP individuales, de los que descarta paquetes con el objetivo de indicar al emisor que debe reducir la tasa de envío.

Weighted Random Early Detection (WRED): WRED es una extensión de la técnica de Random Early Detection (RED), donde una sola cola puede tener varios umbrales de cola diferentes. Cada umbral de cola se asocia a una clase de tráfico particular. Por ejemplo, una cola puede tener umbrales más bajos para los paquetes de prioridad más baja. Una acumulación de cola hará que los paquetes de menor prioridad sean desechados, de esta manera se protege los paquetes

de prioridad más alta en la misma cola. De esta manera la calidad del servicio de priorización se hace posible para paquetes importantes de un grupo de paquetes que usan el mismo búfer. Es más probable que el tráfico normal sea descartado en lugar del tráfico que tiene una mayor prioridad.

CÁLCULO DEL NÚMERO DE LLAMADAS CONCURRENTES PARA ADMINISTRACIÓN DEL ANCHO DE BANDA

El número de llamadas concurrentes resulta ser una estimación, luego de un detallado trabajo de campo, en que se verifica in situ el comportamiento que tiene una organización en función de las llamadas que realizan los personeros de la institución de manera simultánea y el tiempo que demoran en ellas. Una vez que se disponen de estos datos reales, es imprescindible suponer un crecimiento progresivo de dicho número de llamadas, para realizar así los cálculos necesarios de ancho de banda requerido de la red con el propósito de poder administrar a futuro una red telefónica más grande que la actual. En esta forma se asegura un desempeño adecuado de la red de voz.

CÁLCULO DEL ANCHO DE BANDA PARA VoIP (TELEFONÍA IP)

El adecuado cálculo del ancho de banda es uno de los factores más importantes al momento de realizar el diseño de redes de voz sobre protocolo IP, ya que, de esta manera, se asegura la correcta prestación del servicio. El requerimiento del ancho de banda necesario de un enlace, para el transporte de voz IP, está intrínsecamente relacionado

al análisis de dos factores:

Número de llamadas concurrentes: Es la estimación de la máxima cantidad de llamadas que se pueden efectuar de manera simultánea sobre la red. Se debe considerar un margen de crecimiento.

Requerimiento de ancho de banda de cada llamada telefónica: Se debe tener en consideración el códec, opciones de compresión, tipo de enlace que transportará la llamada etc.

De manera simplificada, se utiliza el siguiente método para el cálculo de ancho de banda requerido para paquetes de voz:

Paso 1: Calcular el tamaño de las tramas de voz

El tamaño de la trama está en función del códec utilizado y los encabezados de capa 4, capa 3 y capa 2. Así:

Tamaño de trama=Payload+Enc.4+Enc.3+Enc.2

Para enlaces de bajo ancho de banda, y considerando el tamaño de la trama a transmitir, es aconsejable a veces, utilizar mecanismos de compresión de los encabezados de capa 3 y capa 4, lo que se denomina compresión RTP (cRTP). Con este método, se logra reducir el tamaño de dichos encabezados a 2 o 4 Bytes. Hay que recordar que existen algunos escenarios que no admiten cRTP o que de forma intencional, no se aplica ningún tipo de compresión de encabezados de capa 3 y capa 4. Una vez que se obtiene el valor de la trama en Bytes, es necesario expresar dicha cantidad en bits,

Paso 2: Calcular el ancho de banda requerido por una sola llamada

Para efectuar este paso, es necesario seleccionar un códec de voz a ser utilizado. Los códec más utilizados para digitalización de voz son: G.711, G.726 y G.729, los cuales generan 50 paquetes por segundo (PPS).

Para calcular el ancho de banda requerido para cada llamada, se debe multiplicar el tamaño de cada trama por la cantidad de tramas que se envían por segundo.

Paso 3: Calcular el ancho de banda requerido en la implementación
En este punto se debe considerar el número de llamadas concurrentes y multiplicar, el ancho de banda requerido para una llamada, por el número de llamadas concurrentes que se haya estimado en la red.

De esta manera, se encuentra de forma teórica el Ancho de Banda necesario para el transporte de paquetes de voz IP.

CAPÍTULO V

DEFINICIÓN DE LA PLATAFORMA DEL MODELO DE CALIDAD DE SERVICIO, Y TÉCNICA DE QUEUING A UTILIZAR.

INTRODUCCIÓN

En el presente capítulo, se presenta un sustento teórico que justifica el uso del modelo de calidad de servicio y el método de encolamiento para dicho modelo, que será implementado en el presente trabajo. Dentro de los modelos de calidad de servicios que podemos encontrar, se ha escogido el modelo DiffServ y se explicarán las razones por las que se lo ha seleccionado para ser implementado conjuntamente con el protocolo IPv6. También se hablará sobre el retardo, que es el factor con un impacto serio sobre la calidad de servicio para tráfico de tiempo real. Se explica, por último, cada uno de los factores que se han encontrado que ocasionan dicho retardo, y cómo pueden ser solucionados con el método de encolamiento LLQ, para garantizar calidad de servicio al tráfico RTP.

Elección del modelo de calidad de servicio

Como se mencionó en capítulos anteriores, calidad de servicio es la capacidad de una red para proporcionar diversos niveles de servicio a los diferentes tipos de tráfico a través de diferentes mecanismos denominados "modelos de calidad de servicio", que, a su vez, cada uno incorpora diferentes métodos de implementación para proporcionar estos tipos de niveles. A continuación se describen

algunas generalidades que justifican el uso del modelo DiffServ para el presente trabajo:

- Es el modelo de calidad de servicio más reciente, creado justamente para superar todas las falencias que se tenían con modelos anteriores, modelos que no fueron creados para redes actuales donde existen diferentes flujos de tráfico corriendo por la misma plataforma, incorporando en este la habilidad de diferenciación de tráfico con varios niveles de servicio, lo cual se ajusta perfectamente para llevar a cabo el presente trabajo.
- Es un modelo denominado Soft QoS, lo cual indica que la asignación de recurso será administrado por una persona en base varios criterios de prioridades de los distintos flujos de las aplicaciones corriendo por la misma plataforma, característica necesaria, ya que se tendrá que realizar dicha administración del recurso de acorde a lo que se considere es necesario para priorizar el tráfico RTP.
- Es un modelo que se basa en una pre configuración de los equipos en base a varios criterios de prioridades del tráfico que son determinados previo a un análisis, pre configuración que facilitara la priorización de tráfico RTP frente al resto de tráfico, proceso que puede ser realizado gracias a su habilidad de soporte de varios niveles de servicios, esta pre configuración en los equipos se denomina PHB (Per Hop Behavior), comportamiento de paquetes por saltos, que consiste en pre configurar cada uno de los nodos para que se comporte de una manera específica frente a un paquete determinado.

Adicional a esto, IPv6 ha mejorado sustancialmente en temas de seguridad y calidad de servicio en comparación con su antecesor IPv4. Incorporando en su cabecera campos como Etiqueta de flujo, que

conjuntamente con el modelo DiffServ que permite diferenciación de tráfico con distintos niveles de servicio, se convierte en la elección más sensata al momento de verificar las perfeccionadas características de clasificación y marcado con las que cuenta IPv6, ya que se podrá asignar el mejor nivel de prioridad al tráfico considerado más importante, y mejorar la eficiencia en la entrega de paquetes gracias al campo etiqueta de flujo de la cabecera IPv6 debido a que ya no se tendrá la necesidad de procesar protocolos, IP origen, puerto origen, IP destino, puerto destino, ya que el flujo será identificado directamente por la etiqueta, (Arbili, 2013).

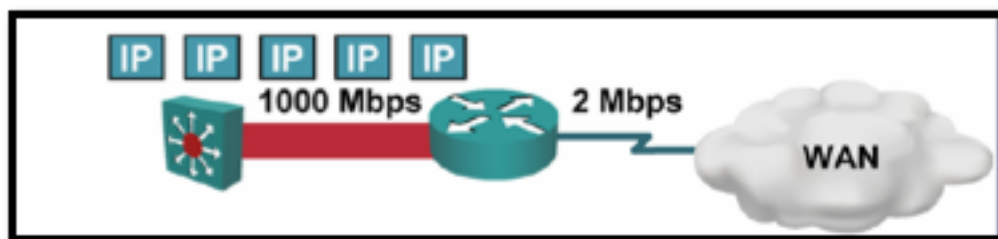
Elección del tipo de encolamiento para el tráfico de voz

A lo largo del desarrollo teórico, se pudieron conocer diversos métodos de encolamiento, entre ellos, LLQ (Low Latency Queuing), el cual ha sido seleccionado para llevar a cabo el siguiente trabajo. A continuación, se describen factores que impiden la entrega inmediata de los paquetes sensible al retardo, una breve descripción del mecanismo de encolamiento CBWFQ que es la base de LLQ, y los motivos que justifican el uso de LLQ para solucionar estos problemas, y así garantizar una entrega inmediata de los paquetes RTP:

- Un factor a tomar en cuenta al momento de querer tener una óptima transmisión de los paquetes a través de los equipos terminales, "Routers", en especial aquellos paquetes que son sensibles a retardo como lo son los paquetes RTP, es el problema de la congestión, que como tal conlleva a la pérdida de paquetes debido a la dificultad de retener los paquetes que llegan a la interfaz de los equipos terminales, ya sea por limitado ancho de banda, que sería el caso del presente trabajo, o falta de tamaño en buffer.

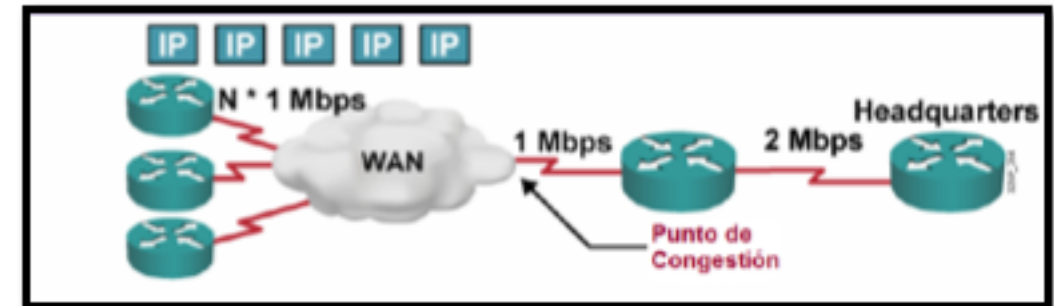
- Una de las causas que ocasiona dicha congestión se denomina “Speed mismatch” (desajuste de velocidad), que se produce debido a que la velocidad de entrada del tráfico en una interfaz, excede la velocidad de la interfaz de salida, como muestra el gráfico 19, que tenemos trabajando un enlace LAN trabajando a 1000 Mbps intentando salir a la WAN por un enlace de 2Mbps, lo cual indica que si todos los usuarios intentan transmitir a máxima velocidad a algún servidor que está en la WAN, lo más probable sería que haya congestión en la interfaz Serial debido a que la velocidad con que entran los paquetes a la interfaz LAN va a ser mucho mayor a la velocidad que intentan salir por la interfaz serial (WAN), causando que los dispositivos intermedios entren en un modo denominado “Tail Drop” (Desecho indiscriminado de paquetes), causando la eliminación de paquetes sin importar de qué tipo de tráfico este sea, es el comportamiento por defecto de los enrutadores, comportamiento que debería ser evitado con mecanismos de encolamiento.

Gráfico 19: Ejemplo del problema Speed Mismatch
(Arbili, 2013).



- Otra posible causa de la congestión es la confluencia y agregación de tráfico, ver gráfico 20.

Gráfico 20: Ejemplo de Confluencia y agregación de tráfico.
(Arbili, 2013).

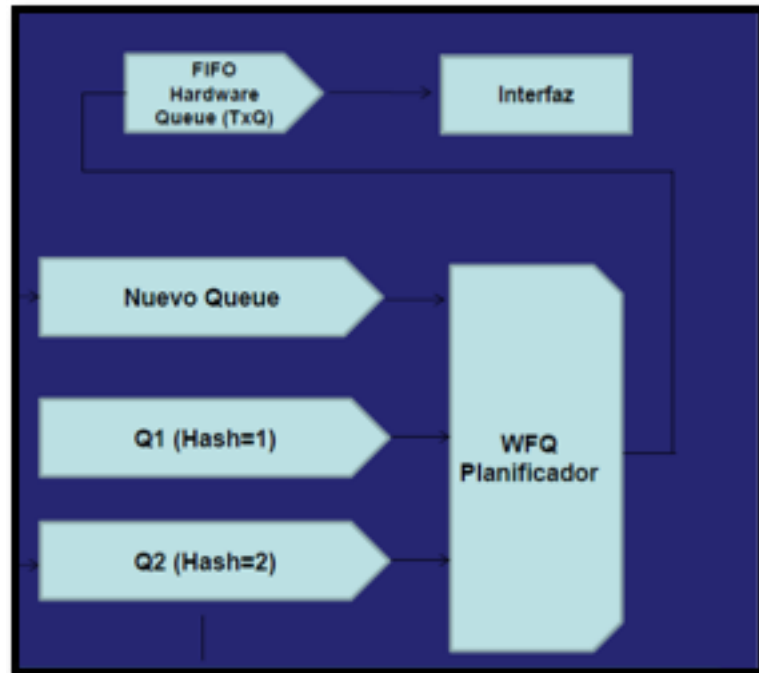


Tal como nos muestra el gráfico, tenemos tres sitios remotos que transmiten por la WAN a 1Mbps cada uno hacia un router de Edge que recibe a 1Mbps, teniendo como punto de congestión el enlace de 1Mbps al momento que los tres sitios remotos quieren transmitir a su máxima velocidad, teniendo un enlace igual a la suma de los tres sitios remotos, que sería 3Mbps, lo cual se tendrían 3Mbps tratando de entrar por un enlace de 1Mbps.

Hasta el momento, nos hemos referido al problema de la congestión y las razones que causa dicho problema. A continuación, se realizará una breve explicación del funcionamiento de un sistema de encolamiento CBWFQ, que es la base del sistema de encolamientos LLQ, para luego mencionar las características de la técnica de encolamiento LLQ por las cuales ha sido seleccionado para resolver estos inconvenientes en el presente trabajo.

Gráfico 21: Funcionamiento de CBWFQ.

(Arbili, 2013).



Como se muestra en el gráfico 21, una cola está conformada por las siguientes partes:

- Una cola en hardware de tipo FIFO que da la cara directamente a la salida de la interfaz.
- Varias colas en software, una por cada clase, (cada clase define una cola en software).
- Un planificador, encargado de despachar los paquetes que están en la cola en software, de acuerdo a las políticas previamente configuradas de calidad de servicio a cada clase.

Teniendo este esquema de funcionamiento de este sistema de encolamiento, lo que sucede al momento que comienza a llegar paquetes hacia una interfaz, es que el dispositivo verifica si hay espacio en la cola en hardware, si es que lo hay, los paquetes pasan directamente

a la cola en hardware, que es la que está directamente a la interfaz de salida por la cual en paquete tendrá un despacho más acelerado, si la cola en hardware está llena, los paquetes se quedarían en cola en software, ubicando los paquetes acorde como se halla clasificado el tráfico, de mayor a menor prioridad, creando una cola por cada prioridad, luego estas colas en software pasarían al planificador, el cual va a planificar el envío de los paquetes para pasarlos a cola en hardware acorde a las políticas de prioridad asignada a cada cola, hay que tomar en cuenta que el planificador se toma su tiempo para realizar dicho proceso, lo cual agrega retardo en este proceso, una vez llegado a cola en hardware, que trabaja de forma FIFO, pasaría directamente a la interfaz de salida.

Habiendo expuesto el funcionamiento del esquema anteriormente mencionado, a continuación, se describirá las características de la técnica de encolamiento LLQ por las cuales ha sido seleccionado para el presente trabajo.

- LLQ puede hacer una clasificación de tráfico personalizada mediante la asignación de clases, teniendo diferentes tipos de tráfico por cada clase creada, pudiendo así, mediante esta característica dividir el tráfico en clases que va a cruzar por la red, que en nuestro caso sería tráfico RTP y de video.
- A partir de la división de tráfico por clases, las clases pueden ser administradas acorde a prioridades, de mayor a menor prioridad cada una de las clases creadas, pudiendo asignar un mínimo ancho de banda a cada clase, garantizando con esto, en situaciones de congestión en la red, un mínimo ancho de banda, dicho ancho de banda puede ser mayor en caso de que no haya congestión en la red,

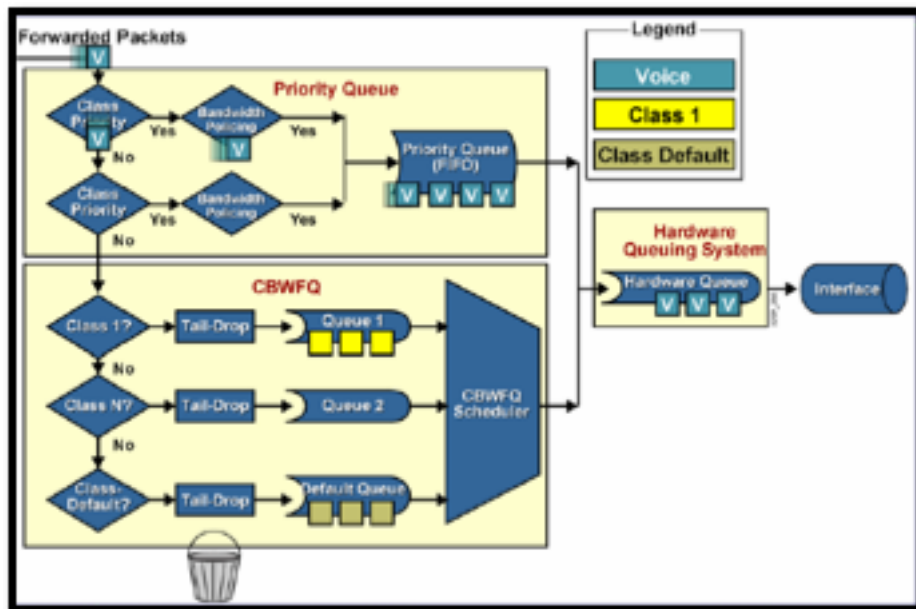
característica que facilita la reserva de un ancho de banda mínimo para la clase a la cual estará asociado el tráfico RTP en este trabajo, que es el tráfico que necesita ser priorizado en cuanto a recurso de la red.

- LLQ es prácticamente un mecanismo CBWFQ con la diferencia que incorpora colas (clases) prioritarias, para tráfico en tiempo real, ver gráfico 22, esto indica que las colas pasan directamente a cola en hardware evitando el retardo del proceso del planificador, teniendo así un despacho acelerado de paquetes, característica exclusiva para tráfico de tiempo real, como es el caso del presente trabajo, donde se requiere minimizar al máximo cualquier tipo de retardo para el tráfico RTP.

Gráfico 22: Arquitectura de LLQ.

(Arbili, 2013).

- Las colas prioritarias que incorpora LLQ, son como se mencionó, destinadas a tráfico de alta prioridad, ofreciendo con estas bajo retardo y reserva de ancho de banda mínimo ideal para tráfico en tiempo real, como lo es el tráfico de VoIP que se maneja en el presente trabajo, (Arbili, 2013)..



BIBLIOGRAFÍA

Bibliografía

3CX (Sin fecha). QoS Quality of service VoIP. United States of America: VoipForo. Recuperado de: <http://www.voipforo.com/QoS/QoS_Jitter.php>

Ahmed, M., T Litchfield, A., Ahmed, S., Mahmood, A., & Hossain , E. (2014). VoIP Performance Analysis over IPv4 and IPv6. Modern Education and computer science PRESS , 43-48.

Álvarez, S., & Gonzalez, A. (2005). Estudio y configuración de calidad de servicio para protocolos ipv4 e ipv6 en una red de fibra óptica WDM. 104-113.

Arbili, Y. E. (13 de Marzo de 2013). Quality of service. Obtenido de <http://www.slideshare.net/>: <http://www.slideshare.net/yasserhh/quality-of-service-45803301>

Cisco (2006). Understanding Delay in Packet Voice Networks. Recuperado de: <http://www.cisco.com/c/en/us/support/docs/voice/voice-quality/5125-delay-details.html>

Cisco (2006). Understanding Delay in Packet Voice Networks. United States of America: Elsevier: Cisco. Recuperado de: <<http://www.cisco.com/c/en/us/support/docs/voice/voice-quality/5125-delay-details.html>>.

Cisco (2007). Performance-Comparison Testing of IPv4 and IPv6 Throughput and Latency on Key Cisco Router Platforms. Disponible en: <http://www.cisco.com/c/dam/en/us/products/collateral/ios-nx-os->

software/enterprise-ipv6-solution/IPv6perf_wp1f.pdf.

Cisco, Voz sobre IP – Consumo de ancho de banda por llamada

Cisco. (Junio de 2016). Instituto Tecnológico de Roque. Obtenido de Instituto Tecnológico de Roque: <http://itroque.edu.mx/cisco/cisco1/course/module6/6.1.3.2/6.1.3.2.html>

Course Hero (2015). Enterprise Readiness for IPv6. Recuperado el 17 de enero de 2016 en <https://www.coursehero.com/file/10834283/Enterprise-Readiness-for-IPv6/>

Elastix (Sin fecha). QoS-Calidad de Servicio para VoIP. United States of America: Elastix. Recuperado de <http://elastixtech.com/qos-calidad-de-servicio-para-voip/>.

España M. Servicios avanzados de telecomunicación (2003).

Flannagan, M. (2001). Administering Cisco QoS for IP Networks (Paperback). Syngress.

Gamboa, F., López, D., & Salcedo, O. (2012). Voice transport ipv4 and ipv6 environments. Visión Electrónica(I).

Gerometta, O. (10 de Febrero de 2015). SlideShere. Obtenido de

Gerometta, O. (2016). Protocolo de Internet versión 6, versión 3.1. Argentina. Edubooks.

IETF. (Junio de 1995). The Internet Engineering Task Force. Obtenido de IETF: <https://www.ietf.org/rfc/rfc1809.txt>

Lenis, E., Eugenio, A., Méndez, A., & Guefri, L. (2015). Análisis de rendimiento en redes IPv6. Entramado , XI (1), 214-229.

Minoli, D. (2006). Voice Over IPv6, Architectures for Next Generation VoIP Networks. United States of America: Elsevier

mIreLucX. (26 de Marzo de 2009). Overblog. Obtenido de Overblog: <http://mirelucx.over-blog.com/article-29483351.html>

Ortega, L. (13 de Febrero de 2012). IPv6. Obtenido de IPv6: <http://ipv6-equipo5.blogspot.com/>

Palacios, A., Salcedo, O., & Lopez, D. (2011). Desempeño de la calidad del servicio (QoS) sobre IPv6. Disponible en <http://www.scielo.org.co/pdf/tecn/v15n28/v15n28a04.pdf>

Rosario, M. A. (2006). El Estándar VoIP - Redes y servicios de banda ancha. Perú: Universidad Nacional Mayor de San Marcos. Recuperado de: <http://www.monografias.com/trabajos33/estandar-voip/estandar-voip.shtml>.

Salcedo, O., Danilo, L., & Rios, A. (2011). Desempeño de la calidad de servicio (QoS) sobre IPv6. Tecnura, XV(28), 32-41.

Shenker, S (1994). Integrated Services in the Internet Architecture: an overview. RFC 1633.

Sistemas y redes digitales. Ginebra: Unión Internacional de Telecomunicaciones.

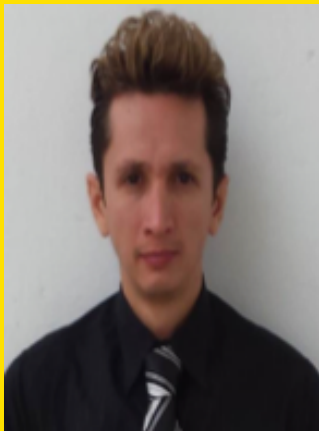
Tutorials Point, Learn IPv6 Internet Protocol Version 6.

UIT-T G.114 (2004). Serie G: Sistemas y medios de transmisión VOIP

Wiki (2016). A reference guide to all things VOIP. Recuperado el 13 de diciembre de 2015 en <http://www.voip-info.org/wiki/view/QoS>.

Unión Internacional de Telecomunicaciones (2003). G.114 : Tiempo de transmisión en un sentido. Recuperado de: <https://www.itu.int/rec/T-REC-G.114-200305-I/es/T-REC-G.114-200305-I!!PDF-S.pdf>

Vesga-Ferreira, J. C., Granados-Acuña, G., & Vesga-Barrera, J. A. (2016). Evaluation of the performance of a network. Iteckne.



Jorge Luis Moreira Calderón, Docente de la Universidad Laica “Eloy Alfaro” de Manabí en su Extensión en Chone, Carrera Ingeniería en Sistemas, con funciones de: Responsable de la comisión de investigación del área técnica, Tutor académico, y docente de la materia de redes de computadoras I y II. Formación profesional: Ingeniero en Informática, Master en redes de comunicaciones. Docente especializado en el campo de las redes de comunicaciones y de forma específica en: Redes LAN, WAM, Satélite, redes celulares, inalámbricas, seguridad y gestión de redes, con visión a las tendencias tecnológicas, dominio de los aspectos técnicos y de la aplicación de la tecnología en el entorno nacional e internacional.



Cristhian Gustavo Minaya Vera, Docente de la Universidad Laica “Eloy Alfaro” de Manabí en su Extensión en Chone Carrera Ingeniería en Sistemas, con funciones de: Responsable de Prácticas y Pasantías de las Carreras Ingeniería en Sistemas, Ingeniería Eléctrica e Ingeniería Civil, Tutor académico, Tutor de Trabajo de Titulación y Miembro de Proyectos de Vinculación. Formación profesional: Licenciado en Informática, Master en Educación Informática. Docente especializado en materias de: Estructura de datos, Programación Orientada a Objetos, Programación Web, Programación Móvil, Herramientas CASE, Sistemas Distribuidos entre otras.



Frank Aquino Cornejo Moreira, Docente del Centro de Aprendizaje y Aplicaciones Informáticas de la Escuela Superior Politécnica Agropecuaria de Manabí “Manuel Félix López”, con funciones de: Impartir cursos a estudiantes de todas las Carreras y al sistema de Admisión y Nivelación de la universidad. Formación profesional: Ingeniero en Informática, Master en Redes de Comunicación. Docente especializado en materias de: Redes de datos, Teleinformática, Estructura de datos, base de datos, Programación de sistemas Operativos, Programación Orientada a Objetos, Programación Web.

ISBN: 978-9942-750-04-4

