

IPv6.br

Curso IPv6 básico

egi.br nic.br

CURSO IPv6 BÁSICO

Rodrigo Regis dos Santos
Antônio M. Moreiras
Eduardo Ascenço Reis
Ailton Soares da Rocha

Núcleo de Informação e Coordenação do ponto BR

São Paulo
2010

Núcleo de Informação e Coordenação do Ponto BR

Diretor Presidente

Demi Getschko

Diretor Administrativo

Ricardo Narchi

Diretor de Serviços

Frederico Neves

Diretor de Projetos Especiais e de Desenvolvimento

Milton Kaoru Kashowakura

Centro de Estudos e Pesquisas em Tecnologia de Redes e Operações – CEPTR0.br

Antônio Marcos Moreiras

Coordenação Executiva e Editorial: Antônio Marcos Moreiras

Autores / Design / Diagramação

Rodrigo Regis dos Santos

Antônio Marcos Moreiras

Eduardo Ascenço Reis

Ailton Soares da Rocha

Versión en Español con el apoyo de Internet Society.

Coordinación: LACNIC

Traducido por: Laureana Pavón Caelles



Latin American and Caribbean Internet Addresses Registry
Registro de Direcciones de Internet para América Latina y Caribe
Registro de Endereços da Internet para América Latina e Caribe

Acerca del Proyecto IPv6.br

IPv6 es la nueva generación del Protocolo de Internet.

El protocolo ya está en uso desde hace algún tiempo, pero ahora es necesario acelerar su despliegue ya que es imprescindible para la continuidad del crecimiento y la evolución de Internet.

El objetivo del proyecto IPv6.br de NIC.br es promover la utilización del nuevo protocolo en Internet y en las redes brasileiras. Para obtener más información diríjase a www.ipv6.br o póngase en contacto escribiendo a ipv6@nic.br.

O **CEPTRO.br** – el Centro de Estudios e Investigaciones en Tecnología de Redes y Operaciones de NIC.br – es responsable de los proyectos que buscan mejorar la calidad de Internet en Brasil y diseminar su utilización, con especial énfasis en sus aspectos técnicos y de infraestructura. Puede obtener más información en el sitio web www.ceptro.br.

Acerca de los autores

Rodrigo Regis dos Santos se graduó en Ciencias de la Computación en la Universidad Presbiteriana Mackenzie y actualmente se desempeña como Analista de Proyectos en NIC.br. Rodrigo es especialista en IPv6 y uno de los responsables del proyecto IPv6.br, cuyo objetivo es incentivar el uso del protocolo en el país.

Antônio M. Moreiras se desempeña en NIC.br coordinando los proyectos relacionados con la infraestructura de Internet y su desarrollo en Brasil. Actualmente está trabajando en la diseminación de IPv6 en Brasil, sincronización de la Hora Legal por parte de NTP, mediciones de la calidad de Internet en Brasil y estudios sobre la web brasileira. Ingeniero Electricista con una Maestría en Ingeniería de la Escuela Politécnica de la Universidad de São Paulo y un MBA en Gestión Empresarial de la Facultad de Administración y Ciencias Contables de la Universidad Federal de Río de Janeiro, actualmente está cursando una especialización en Gobernanza de Internet en la Diplo Foundation.

Eduardo Ascenço Reis es especialista en redes IP, sistemas Unix y servicios de Internet. Como formación, tiene Especialización en Redes de Computadoras por LARC/USP y Licenciado en Ciencias Biológicas por la USP. Su experiencia profesional en TI es desde 1995, trabajando en las empresas: Universidad de São Paulo (USP), Ericsson Brasil, comDominio (IDC – AS16397) y CTBC Multimedia (NSP, ISP - AS27664). Actualmente actúa como supervisor de proyecto en el Centro de Estudios e Investigación en Tecnología de Redes y Operaciones (CEPTRO.br) en el Centro de Información y Coordinación del Punto BR (NIC.br) y es uno de los responsables por el PTTMetro (PTT.br).

Ailton Soares da Rocha es Analista de Proyectos en NIC.br, donde trabaja en investigaciones y proyectos relacionados con la infraestructura de Internet en el país y ha tenido una destacada actuación en los proyectos PTT.br, IPv6.br y NTP.br. Ingeniero Electricista y en Comunicaciones graduado del Instituto Nacional de Telecomunicaciones (INATEL), desde hace más de 10 años trabaja en la coordinación del área de redes e Internet de la institución.



IPv6.br



IPv6.br

Acerca de la licencia



Atribuição-Compartilhamento pela mesma Licença 2.5 Brasil

Você pode:



copiar, distribuir, exhibir e executar a obra



criar obras derivadas



Sob as seguintes condições:



Atribuição. Você deve dar crédito ao autor original, da forma especificada pelo autor ou licenciante.



Compartilhamento pela mesma Licença. Se você alterar, transformar, ou criar outra obra com base nesta, você somente poderá distribuir a obra resultante sob uma licença idêntica a esta.

- Los términos de la licencia de esta obra se deberán aclarar para cada nuevo uso o distribución.
- En caso de creación de obras derivadas, no se deberán utilizar los logotipos de CGI.br, NIC.br, IPv6.br y CEPTRO.br.
- Al atribuir su autoría, esta obra debe ser citada de la siguiente manera:
 - Apostilla "Curso IPv6 Básico" de NIC.br, disponible en el sitio <http://curso.ipv6.br> o escribiendo a ipv6@nic.br.
- Cualquiera de estas condiciones puede no aplicarse si se obtiene el permiso del autor.
En caso de ser necesario, NIC.br puede ser consultado escribiendo a ipv6@nic.br.
- Ninguna parte de esta licencia perjudica o restringe los derechos morales del autor.

IPu6.br

La nueva generación del
Protocolo de Internet

Introducción

Módulo 1

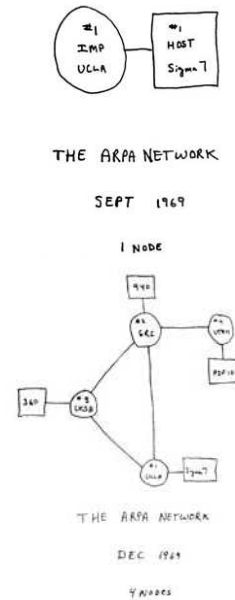
En este módulo introductorio comenzaremos por aprender algo de la historia de Internet y del desarrollo del protocolo IP.

Comprenderemos cuáles son los problemas que provoca la forma inicialmente adoptada para distribuir las direcciones IP y el rápido crecimiento de Internet, y cuáles son las soluciones adoptadas para resolver estos problemas. Luego de este repaso histórico veremos cómo algunas de estas soluciones evolucionaron hasta llegar a la versión 6 del protocolo IP, conocido como IPv6.

En este módulo también veremos, a través de datos estadísticos, la necesidad de desplegar IPv6 en las redes de computadoras, comparando datos sobre la tasa de crecimiento de Internet y la adopción y utilización de IPv6. También discutiremos las consecuencias de no implementar el nuevo protocolo IP y del uso generalizado de técnicas consideradas paliativas, como por ejemplo las NAT.

Internet y TCP/IP

- 1969 – Inicio de ARPANET
- 1981 – Definición de IPv4 en la RFC 791
- 1983 – ARPANET adopta TCP/IP
- 1990 – Primeros estudios sobre el agotamiento de las direcciones
- 1993 – Internet comienza a ser explotada comercialmente
 - Se intensifica el debate sobre el posible agotamiento de las direcciones libres y el crecimiento de la tabla de enrutamiento.



8

En 1966 el Departamento de Defensa (DoD - *Department of Defense*) del gobierno de Estados Unidos, a través de su Agencia de Investigación de Proyectos Avanzados (ARPA - *Advanced Research Projects Agency*), inició un proyecto para interconectar las computadoras de los centros militares y de investigación. Este sistema de comunicación y control distribuido con fines militares recibió el nombre de ARPANET; su principal objetivo teórico era crear una arquitectura de red sólida y robusta que, incluso en caso de falla de alguna estación, pudiera trabajar con las computadoras y enlaces restantes. En 1969 se instalaron los primeros cuatro nodos de la red en la Universidad de Los Ángeles (UCLA), la Universidad de California en Santa Bárbara (UCSB), el Instituto de Investigaciones de Standford (SRI) y la Universidad de Utah.

En sus inicios, ARPANET trabajaba con diferentes protocolos de comunicación, siendo NCP (*Network Control Protocol*) el principal. El 1/1/1983, cuando la red ya contaba con 562 hosts, todas las máquinas de ARPANET adoptaron como estándar los protocolos TCP/IP, permitiendo así el crecimiento ordenado de la red y eliminando las restricciones que presentaban los protocolos anteriores.

Definido en la RFC 791, el protocolo IP tiene dos funciones básicas: la fragmentación, que permite enviar paquetes mayores que el límite de tráfico establecido para un enlace dividiéndolos en paquetes más pequeños, y el direccionamiento, que permite identificar el destino y el origen de los paquetes en base a la dirección almacenada en el encabezado del protocolo. La versión del protocolo IP que se utilizaba en aquella época y continúa utilizándose en la actualidad es la versión 4 o IPv4. Esta versión demostró ser muy robusta, de fácil implementación e interoperabilidad, no obstante lo cual el proyecto original no previó algunos aspectos tales como:

- El crecimiento de las redes y el posible agotamiento de las direcciones IP;
- El crecimiento de la tabla de enrutamiento;
- Problemas relacionados con la seguridad de los datos transmitidos;
- Prioridad en la entrega de determinados tipos de paquetes.

Agotamiento de las direcciones IPv4

- IPv4 = 4.294.967.296 direcciones.
- Política inicial de distribución de direcciones.

- Clase A
 - IBM
 - HP
 - AT&T
 - MIT
 - DoD
 - US Army
 - USPS
 -
- Clase B
- Clase C
- Direcciones reservadas

9

Las especificaciones de IPv4 reservan 32 bits para direccionamiento, permitiendo generar más de 4 mil millones de direcciones diferentes. Inicialmente, estas direcciones se dividieron en tres clases de tamaño fijo de la siguiente manera:

- **Clase A:** Definía el bit más significativo como 0, utilizaba los 7 bits restantes del primer octeto para identificar la red y los 24 bits restantes para identificar el *host*. Estas direcciones utilizaban el rango de **1.0.0.0 a 126.0.0.0**;
- **Clase B:** Definía los 2 bits más significativos como 10, utilizaba los 14 bits siguientes para identificar la red y los 16 bits restantes para identificar el *host*. Estas direcciones utilizaban el rango de **128.1.0.0 a 191.254.0.0**;
- **Clase C:** Definía los 3 bits más significativos como 110, utilizaba los 21 bits siguientes para identificar la red y los 8 bits restantes para identificar el *host*. Estas direcciones utilizaban el rango de **192.0.1.0 a 223.255.254.0**;

Clase	Formato	Redes	Hosts
A	7 Bits Rede, 24 Bits Host	128	16.777.216
B	14 Bits Rede, 16 Bits Host	16.384	65.536
C	21 Bits Rede, 8 Bits Host	2.562.097.152	256

Aunque la intención de esta división era flexibilizar la distribución de direcciones abarcando redes de diferentes tamaños, este tipo de clasificación demostró ser ineficiente. Así, la clase A atendía un número muy pequeño de redes pero ocupaba la mitad de todas las direcciones disponibles; para direccionar 300 dispositivos en una red era necesario obtener un bloque de direcciones clase B, con lo cual se desperdiciaba prácticamente la totalidad de las 65 mil direcciones; y las 256 direcciones clase C no satisfacían las necesidades de la gran mayoría de las redes.

Otro factor que contribuyó al desperdicio de direcciones fue el hecho de que decenas de rangos clase A fueron asignados íntegramente a grandes organizaciones tales como IBM, AT&T, Xerox, HP, Apple, MIT, Ford, el Departamento de Defensa de Estados Unidos, entre muchas otras, poniendo a disposición de cada una de ellas 16.777.216 millones de direcciones. Además, se reservaron 35 rangos de direcciones clase A para usos específicos tales como *multicast*, *loopback* y uso futuro.

En 1990 ya había 313.000 *hosts* conectados a la red y algunos estudios comenzaban a hablar de un colapso provocado por la falta de direcciones. Otros problemas tales como el crecimiento de la tabla de enrutamiento también se agudizaban a medida que Internet evolucionaba.

Debido a la tasa de crecimiento de Internet y a la política de distribución de direcciones, en mayo de 1992 ya se habían distribuido 38% de los rangos de direcciones clase A, 43% de la clase B y 2% de la clase C. En esa época ya había 1.136.000 *hosts* conectados a la red.

En 1993, la creación del protocolo HTTP y la liberación de Internet por parte del Gobierno de Estados Unidos para su utilización comercial produjeron un salto aun mayor en la tasa de crecimiento de la red, que pasó de 2.056.000 de *hosts* en 1993 a más de 26.000.000 de *hosts* en 1997.

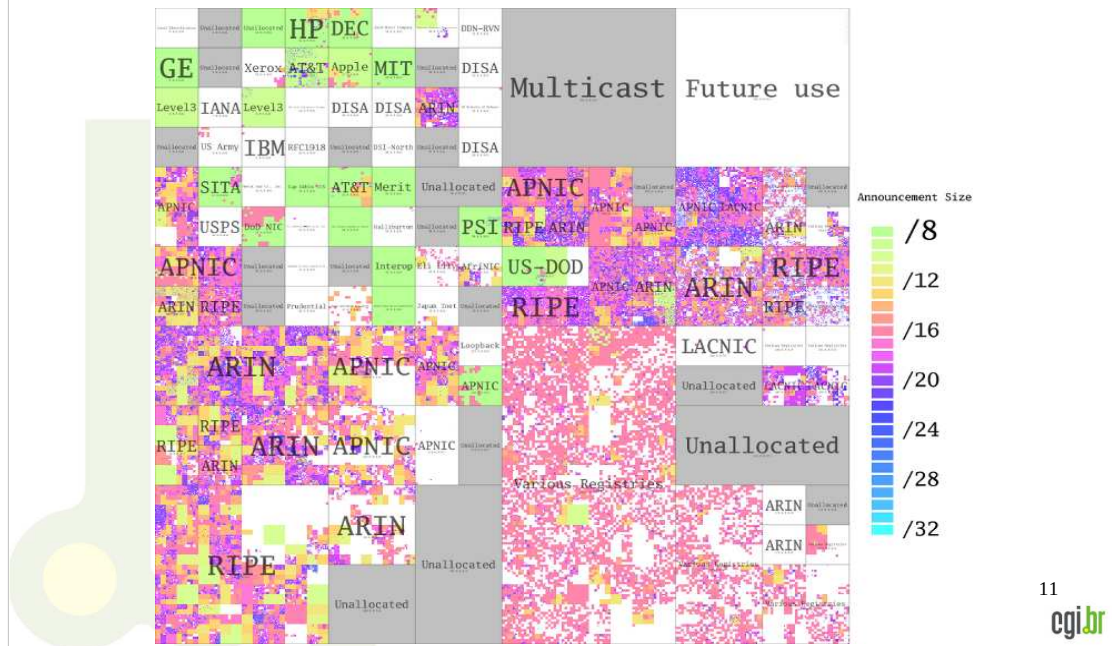
Fecha	Hosts	Dominios
1981	213	-
1982	235	-
1983	562	-
1984	1.024	-
1985	1.961	-
1986	5.089	-
1987	28.174	-
1988	56.000	1.280
1989	159.000	4.800
1990	313.000	9.300
1991	617.000	18.000
1992	1.136.000	17.000
1993	2.056.000	26.000
1994	3.212.000	46.000
1995	8.200.000	120.000
1996	16.729.000	488.000
1997	26.053.000	1.301.000

Tabla de crecimiento de Internet.

Más información:

- RFC 1287 - *Towards the Future Internet Architecture*.
- RFC 1296 - *Internet Growth (1981-1991)*
- Solensky F., 'Continued Internet Growth', *Proceedings of the 18th Internet Engineering Task Force*, Agosto 1990, <http://www.ietf.org/proceedings/prior29/IETF18.pdf>
- IANA IPv4 Address Space Registry - <http://www.iana.org/assignments/ipv4-address-space/ipv4-address-space.xml>
- RFC 3330 - *Special-Use IPv4 Addresses*.

Agotamiento de las direcciones IPv4



Esta imagen muestra la información de la tabla de enrutamiento BGP tomada del proyecto Routeviews. Aquí el espacio de direcciones IPv4 unidimensional se representa en una imagen bidimensional en la cual los bloques CIDR siempre aparecen como cuadrados o rectángulos.

Más información:

- <http://maps.measurement-factory.com/>

Soluciones

Soluciones paliativas:

- 1992 – La IETF crea el grupo ROAD (*ROuting and ADdressing*).
 - CIDR (RFC 4632)
 - Fin del uso de clases = bloques de tamaño apropiado.
 - Dirección de red = prefijo/longitud.
 - Agregación de rutas = reduce el tamaño de la tabla de rutas.
 - DHCP
 - Distribución de direcciones dinámicas.
 - NAT + RFC 1918
 - Permite conectar toda una red de computadoras usando una sola dirección válida en Internet, pero con algunas restricciones.

12

Ante este escenario la IETF (*Internet Engineering Task Force*) comenzó a debatir estrategias para resolver el tema del agotamiento de las direcciones IP y el problema del crecimiento de la tabla de enrutamiento. Para ello, en noviembre de 1991 se formó el grupo de trabajo ROAD (*ROuting and Addressing*), que para solucionar estos problemas propuso la utilización de CIDR (*Classless Inter-domain Routing*).

Definido en la RFC 4632 (que dejó obsoleta la RFC 1519), las ideas básicas detrás del CIDR eran poner fin al uso de clases de direcciones para permitir la distribución de bloques de tamaño adecuado a las necesidades reales de cada red, y la agregación de rutas para reducir el tamaño de la tabla de enrutamiento. Los bloques CIDR se identifican mediante prefijos de red. Por ejemplo, en la dirección **a.b.c.d/x**, los x bits más significativos indican el prefijo de red. Otra manera de indicar el prefijo es a través de máscaras, donde la máscara **255.0.0.0** indica un prefijo /8, **255.255.0.0** indica un /16, y así sucesivamente.

Otra solución, presentada en la RFC 2131 (que dejó obsoleta la RFC 1541), fue el protocolo DHCP (*Dynamic Host Configuration Protocol*). A través del protocolo DHCP un *host* puede obtener una dirección IP automáticamente y adquirir información adicional como por ejemplo la máscara de subred, la dirección del router por defecto y la dirección del servidor DNS local.

DHCP ha sido muy utilizado por los ISP debido a que les permite asignar direcciones IP temporarias a sus clientes conectados. De esta forma ya no es necesario obtener una dirección para cada cliente, sino que alcanza con asignar direcciones dinámicamente a través del servidor DHCP. Este servidor tendrá una lista de direcciones IP disponibles: cada vez que un nuevo cliente se conecte a la red le será asignada una de estas direcciones de forma aleatoria, dirección que será devuelta en el momento que el cliente se desconecte.

Soluciones

• NAT

• Ventajas:

- Reduce la necesidad de direcciones públicas;
- Facilita la numeración interna de las redes;
- Oculta la topología de las redes;
- Solo permite la entrada de paquetes generados en respuesta a un pedido de la red.

• Desventajas:

- Rompe con el modelo end-to-end de Internet;
- Dificulta el funcionamiento de una serie de aplicaciones;
- No es escalable;
- Aumento del procesamiento en el dispositivo traductor;
- Falsa sensación de seguridad;
- Imposibilidad de rastrear el camino del paquete;
- Imposibilita el uso de algunas técnicas de seguridad tales como IPSec.

13

cgi.br

NAT (*Network Address Translation*) fue otra técnica paliativa desarrollada para resolver el problema del agotamiento de las direcciones IPv4. Definida en la RFC 3022 (que dejó obsoleta la RFC 1631), la idea básica es permitir que varios *hosts* puedan salir a Internet con una única dirección IP o con un número pequeño de direcciones IP. Dentro de una red, cada equipo recibe una dirección IP privada única que es utilizada para enrutar el tráfico interno. Sin embargo, cuando un paquete debe ser enrutado fuera de la red, las direcciones IP privadas se traducen a direcciones IP públicas globalmente únicas.

Para que este esquema sea posible se utilizan los tres rangos de direcciones IP declarados como privados en la RFC 1918, siendo la única regla de utilización que ningún paquete que contiene estas direcciones puede atravesar la Internet pública. Los tres rangos reservados son los siguientes:

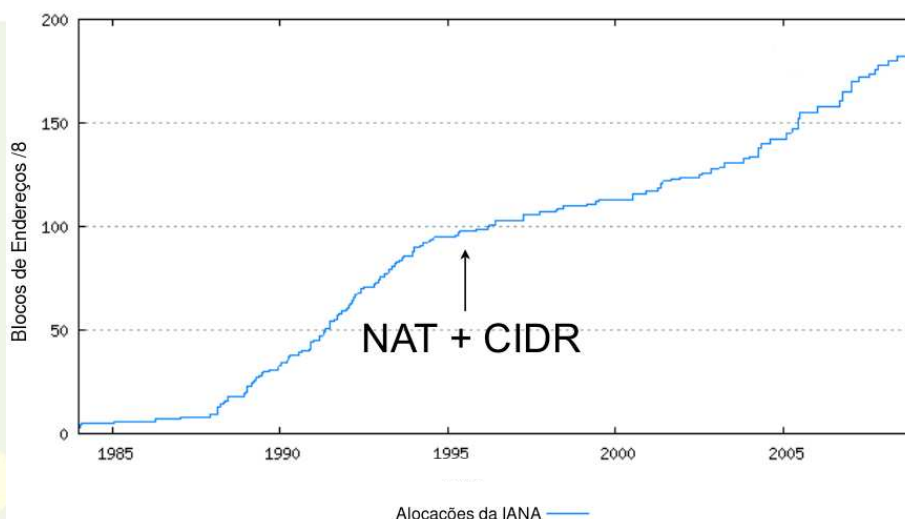
10.0.0.0	a	10.255.255.255 /8	(16.777.216 <i>hosts</i>)
172.16.0.0	a	172.31.255.255 /12	(1.048.576 <i>hosts</i>)
192.168.0.0	a	192.168.255.255 /16	(65.536 <i>hosts</i>)

El uso de NAT demostró ser eficiente en cuanto a la economía de direcciones IP, además de presentar algunos otros aspectos positivos tales como facilitar la numeración interna de las redes, ocultar la topología de las redes y solo permitir la entrada de paquetes generados en respuesta a una solicitud de la red. No obstante, el uso de NAT presenta inconvenientes que no compensan las ventajas que ofrece.

NAT rompe con el modelo end-to-end de Internet ya que no permite conexiones directas entre dos *hosts*, lo que dificulta el funcionamiento de una serie de aplicaciones tales como P2P, VoIP y VPN. Otro problema es la baja escalabilidad, ya que el número de conexiones simultáneas es limitado y además requiere una gran capacidad de procesamiento por parte del dispositivo traductor. El uso de NAT también imposibilita rastrear el camino del paquete mediante herramientas como *traceroute*, por ejemplo, y dificulta la utilización de algunas técnicas de seguridad como IPSec. Además, su uso genera una falsa sensación de seguridad ya que, a pesar de no permitir la entrada de paquetes no autorizados, NAT no realiza ningún tipo de filtrado ni verificación de los paquetes que lo atraviesan.

Soluciones

Soluciones paliativas: Queda apenas 14%



14

cgi.br

Aunque estas soluciones han disminuido la demanda de direcciones IP, no han sido suficientes para resolver los problemas derivados del crecimiento de Internet. La adopción de estas técnicas redujo la cantidad de bloques de direcciones solicitados a la IANA apenas en un 14%, mientras que la curva de crecimiento de Internet continuaba mostrando un aumento exponencial.

De hecho, estas medidas permitieron que hubiera más tiempo para desarrollar una nueva versión del protocolo IP, una versión que se basara en los principios que contribuyeron al éxito de IPv4 pero que también fuese capaz de superar las fallas que se detectaron.

Más información:

- RFC 1380 - *IESG Deliberations on Routing and Addressing*
- RFC 1918 - *Address Allocation for Private Internets*
- RFC 2131 - *Dynamic Host Configuration Protocol*
- RFC 2775 - *Internet Transparency*
- RFC 2993 - *Architectural Implications of NAT*
- RFC 3022 - *Traditional IP Network Address Translator (Traditional NAT)*
- RFC 3027 - *Protocol Complications with the IP Network Address Translator*
- RFC 4632 - *Classless Inter-domain Routing (CIDR): The Internet Address Assignment and Aggregation Plan.*

Soluciones

Estas medidas han permitido contar con más tiempo para desarrollar una nueva versión del protocolo IP.

- 1992 – La IETF crea el grupo IPng (*IP Next Generation*)
 - Aspectos principales:
 - Escalabilidad;
 - Seguridad;
 - Configuración y administración de redes;
 - Soporte para QoS;
 - Movilidad;
 - Políticas de enrutamiento;
 - Transición.

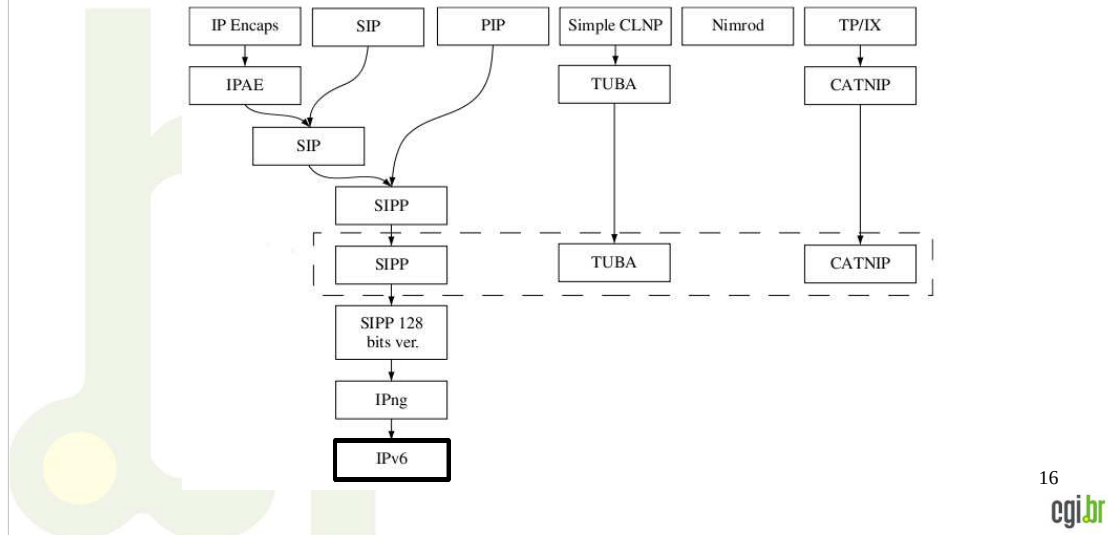
15

Fue así que en diciembre de 1993, a través de la RFC 1550, la IETF formalizó las investigaciones sobre la nueva versión del protocolo IP, solicitando el envío de proyectos y propuestas para el nuevo protocolo. Esta fue una de las primeras acciones del grupo de trabajo de la IETF denominado *Internet Protocol next generation* (IPng). Los principales aspectos que debían ser abordados al elaborar la siguiente versión del protocolo IP eran:

- Escalabilidad;
- Seguridad;
- Configuración y administración de redes;
- Soporte para QoS;
- Movilidad;
- Políticas de enrutamiento;
- Transición.

Soluciones

Solución definitiva:



Varios proyectos comenzaron a estudiar los efectos del crecimiento de Internet, entre los que se destacaron CNAT, *IP Encaps*, *Nimrod* y *Simple CLNP*. De estas propuestas surgieron el TCP and UDP with *Bigger Addresses* (TUBA), que fue una evolución del *Simple CLNP*, y el *IP Address Encapsulation* (IPAE), una evolución del *IP Encaps*. Algunos meses después se presentaron los proyectos *Paul's Internet Protocol* (PIP), *Simple Internet Protocol* (SIP) y TP/IX. Una nueva versión del protocolo SIP, que englobaba algunas funcionalidades de IPAE fue presentada poco antes de agregarse al PIP, obteniéndose como resultado el *Simple Internet Protocol Plus* (SIPP). En este mismo período TP/IX pasó a llamarse *Common Architecture for the Internet* (CATNIP).

En enero de 1995, el IPng presentó un resumen de la evaluación de las tres principales propuestas en la RFC 1752:

- **CATNIP** – Fue concebido como un protocolo de convergencia para permitir que cualquier protocolo de la capa de transporte se ejecutara sobre cualquier protocolo de la capa de red, creando así un ambiente común entre los protocolos de Internet, OSI y Novell;
- **TUBA** – Proponía aumentar el espacio para direccionamiento IPv4 y volverlo más jerárquico, intentando evitar la necesidad de modificar los protocolos de la capa de transporte y aplicación. Pretendía una migración simple y a largo plazo, basada en la actualización de los *host* y servidores DNS, pero sin necesidad de encapsulado o traducción de paquetes ni mapeo de direcciones;
- **SIPP** – Concebido para ser una etapa evolutiva del protocolo IPv4, sin cambios radicales y conservando la interoperabilidad con la versión 4 del protocolo IP, proveía una plataforma para nuevas funcionalidades de Internet, aumentaba el espacio para direccionamiento de 32 bits a 64 bits, presentaba un mayor nivel de jerarquía y estaba compuesto por un mecanismo que permitía “ampliar la dirección” denominado *cluster addresses*. Ya tenía encabezados de extensión y un campo *flow* para identificar el tipo de flujo de cada paquete.

Sin embargo, también según lo informado en la RFC 1752, las tres propuestas también presentaban problemas significativos. Finalmente, la recomendación para el nuevo Protocolo de Internet se basó en una versión revisada del SIPP, que pasó a incorporar direcciones de 128 bits, junto con los elementos de transición y autoconfiguración del TUBA, el direccionamiento basado en el CIDR y los encabezados de extensión. Por ser considerado muy incompleto, el protocolo CATNIP fue descartado.

A raíz de esta definición, la nueva versión del Protocolo de Internet pasó oficialmente a llamarse IPv6.

Más información:

- RFC 1550 - *IP: Next Generation (IPng) White Paper Solicitation*
- RFC 1752 - *The Recommendation for the IP Next Generation Protocol*

IPv6

- 1998 - Definido por la RFC 2460
- 128 bits para direccionamiento.
- Encabezado base simplificado.
- Encabezados de extensión.
- Identificación de flujo de datos (QoS).
- Mecanismos de IPSec incorporados al protocolo.
- Realiza la fragmentación y ensamble de los paquetes solamente en el origen y el destino.
- No requiere el uso de NAT, por lo que permite conexiones end-to-end.
- Mecanismos que facilitan la configuración de redes.
-

18

Las especificaciones de IPv6 fueron inicialmente presentadas en la RFC 1883 de diciembre de 1995, pero en diciembre de 1998 esta RFC fue reemplazada por la RFC 2460. Entre los principales cambios respecto a IPv4 se destacan:

- **Mayor capacidad de direccionamiento:** en IPv6 el espacio de direccionamiento aumentó de 32 bits a 128 bits, permitiendo: niveles más específicos de agregación de direcciones, identificar una cantidad mucho mayor de dispositivos en la red e implementar mecanismos de autoconfiguración. También se mejoró la escalabilidad del enrutamiento *multicast* mediante la adición del campo "alcance" en la dirección *multicast*. También se definió un nuevo tipo de direcciones, las direcciones *anycast*;
- **Simplificación del formato del encabezado:** con el objetivo de reducir el costo de procesamiento de los paquetes en los routers, algunos campos del encabezado IPv4 se eliminaron o se convirtieron en opcionales;
- **Soporte para encabezados de extensión:** las opciones dejaron de formar parte del encabezado base, permitiendo un enrutamiento más eficaz, límites menos rigurosos en cuanto al tamaño y la cantidad de opciones, y una mayor flexibilidad para la introducción de nuevas opciones en el futuro;
- **Capacidad de identificar flujos de datos:** se agregó un nuevo recurso que permite identificar paquetes que pertenecen a determinados flujos de tráfico que pueden requerir tratamientos especiales;
- **Soporte para autenticación y privacidad:** se especificaron encabezados de extensión capaces de proveer mecanismos de autenticación y garantizar la integridad y confidencialidad de los datos transmitidos.

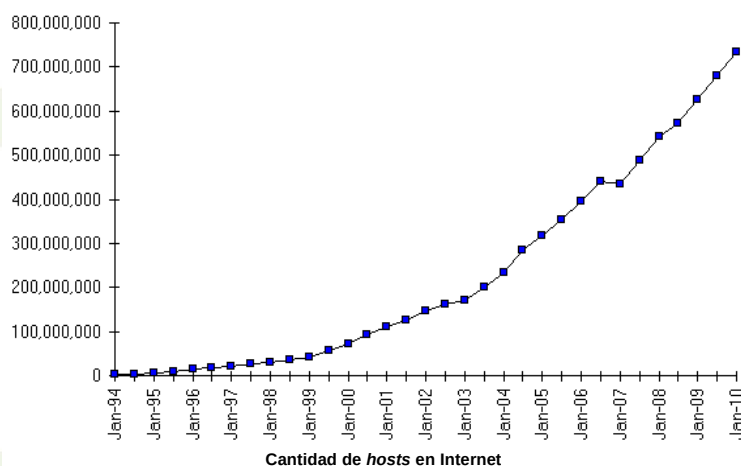
Además, IPv6 también modificó el tratamiento de la fragmentación de los paquetes que pasó a ser realizada solo en el origen; permitió el uso de conexiones end-to-end, principio que se había roto con IPv4 debido al uso generalizado de NAT; aportó recursos que facilitan la configuración de redes, entre otros aspectos que fueron mejorados en relación con IPv4.

Más información:

- RFC 2460 - *Internet Protocol, Version 6 (IPv6) Specification*

¿Por qué utilizar IPv6 hoy?

- Internet continúa creciendo



20

egi.br

Durante los últimos diez años, durante el desarrollo de IPv6, Internet continuó mostrando una tasa de crecimiento cada vez más acelerada. El número de *hosts* conectados a Internet saltó de 30.000.000 a aproximadamente 732.000.000 en la actualidad, con un número cada vez mayor de usuarios y dispositivos conectados a la Red.

¿Por qué utilizar IPv6 hoy?

- Internet continúa creciendo
 - Mundo
 - 1.733.993.741 usuarios de Internet;
 - 25,6% de la población;
 - Crecimiento de 380,3% en los últimos 9 años.
 - Se anticipa que en 2014 el total de celulares, smartphones, netbooks y módems 3G llegará a 2,25 mil millones de aparatos.
 - Brasil
 - 21% de los hogares tienen acceso a Internet;
 - 2,6 millones de conexiones en banda ancha móvil;
 - 11 millones de conexiones en banda ancha fija;

21

Esta expansión de Internet se puede medir a través de diversos factores y numerosas investigaciones han demostrado que este crecimiento no ocurre de manera aislada.

Se estima que en el mundo existen 1.733.993.741 usuarios de Internet (25,6% de la población mundial), lo cual, si consideramos los últimos nueve años, representa un crecimiento de 380,3%. Si esta tasa de crecimiento se mantiene, dentro de dos años habrá 2000 millones de usuarios, superando así el pronóstico de que este número solo se alcanzaría en 2015. La siguiente tabla detalla estas cifras, mostrando la penetración y el crecimiento de Internet en cada región alrededor del mundo.

De acuerdo con datos de ABI Research, se anticipa que la cantidad de equipos móviles con capacidad para acceder a Internet – celulares, smartphones, netbooks y módems 3G – llegará a 2,25 mil millones de aparatos.

<i>Regiões</i>	<i>População (em 2009)</i>	<i>Usuários de Internet (em 2000)</i>	<i>Usuários de Internet (atualmente)</i>	<i>% por Região</i>	<i>% no mundo</i>	<i>Crescimento 2000-2009</i>
África	991.002.342	4.514.400	67.371.700	6,8 %	3,9%	1.392,4 %
Ásia	3.808.070.503	114.304.000	728.257.230	19,4 %	42,6 %	545,9 %
Europa	803.850.858	105.096.093	418.029.796	52,0 %	24,1 %	297,8 %
Oriente Médio	202.687.005	3.284.800	57.425.046	28,3 %	3,3 %	1.648,2 %
América Norte	340.831.831	108.096.800	252.908.000	74,2 %	14,6 %	134,0 %
América Latina /Caribe	586.662.468	18.068.919	179.031.479	30,5 %	10,3 %	890,8 %
Oceania	34.700.201	7.620.480	20.970.419	60,4 %	1,2 %	175,2 %
TOTAL	6.767.805.208	360.985.492	1.733.993.741	25,6 %	100 %	380,3 %

Datos del 15/01/2010.

Siguiendo esta tendencia, en Brasil el porcentaje de hogares con acceso a Internet a través de computadoras personales aumentó de 11% en el segundo semestre de 2005 al porcentaje actual de 21%. En junio de 2009 Brasil llegó a 2,6 millones de conexiones en banda ancha móvil, con un crecimiento de 34,2% en un semestre. El número de conexiones a través de banda ancha fija ha alcanzado un total de 11 millones.

¿Por qué utilizar IPv6 hoy?

- Esto implica que la demanda de direcciones IPv4 también crece:
 - En 2010 ya se asignaron 12 bloques /8 a los RIR;
 - Quedan apenas 14 bloques /8 libres en la IANA, lo que equivale a 5% del total;
 - Los pronósticos actuales señalan que estos bloques se agotarán en 2011;
 - El stock de los RIR debería durar 2 o 3 años más.



23

egi.br

Como consecuencia de este crecimiento, la demanda de direcciones IP también crece considerablemente. En 2010 IANA ya asignó 12 bloques /8 a los RIR con lo que, de los 256 /8 posibles, en este momento solamente quedan 14 bloques sin asignar, es decir 5% del total. Este índice refuerza el pronóstico de 2011 como fecha del agotamiento de sus bloques de direcciones IPv4, principalmente porque el número de solicitudes de bloques de direcciones aumenta año a año.

En setiembre de 2008 los RIR llegaron a un acuerdo sobre la política que será adoptada por la IANA cuando su reserva de direcciones llegue al límite de 5 bloques /8. Estos últimos bloques /8 serán inmediatamente asignados a cada RIR, quienes los distribuirán entre los ISP y Registros Nacionales. Si el stock de la IANA llega al 2011 como se ha pronosticado, el número de direcciones IPv4 disponibles solo se agotará cuando también se agoten las reservas regionales, lo cual podría ocurrir dos o tres años después de la finalización de las reservas de la IANA.

Aunque Internet es conocida por ser una red global sin ninguna coordinación central, existen ciertos órganos responsables por administrar los principales aspectos técnicos necesarios para su funcionamiento. Las atribuciones de estos órganos están distribuidas a nivel mundial, respetando una estructura jerárquica:

La IANA (*Internet Assigned Numbers Authority*) es responsable por la coordinación y distribución global del espacio de direcciones IP y de los ASN. Desde 1998, la IANA pasó a operar a través de ICANN (*Internet Corporation for Assigned Names and Numbers*), una organización internacional dirigida por un consejo de directores provenientes de diversos países que supervisa a ICANN y a las políticas elaboradas por esta organización, trabajando en asociación con los gobiernos, empresas y organizaciones creadas mediante tratados internacionales que participan en la construcción y mantenimiento de Internet.

La responsabilidad por la distribución local de las direcciones IP está a cargo de los Registros de Internet (*Internet Registry - IR*), los cuales se clasifican de acuerdo con su función principal y el alcance territorial dentro de la estructura jerárquica Internet.

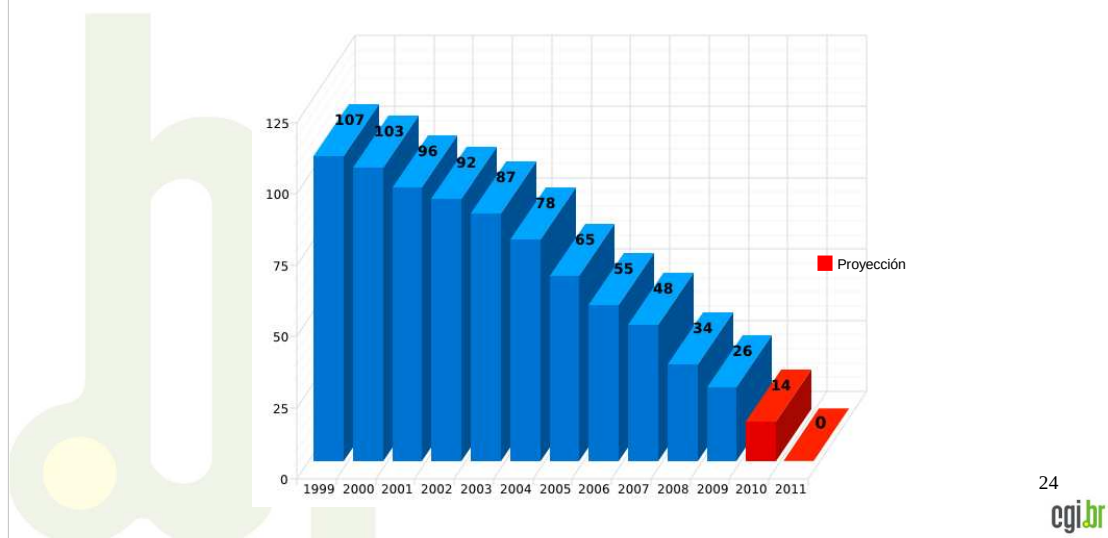
Los Registros Regionales de Internet (*Regional Internet Registries - RIR*) son organismos reconocidos por la IANA que actúan y representan grandes regiones geográficas. La función fundamental de un RIR es administrar y distribuir direcciones IP públicas dentro de sus respectivas regiones. Existen cinco RIR: el *African Network Information Centre (AfriNIC)*, que actúa en la región de África; el *Asia-Pacific Network Information Centre (APNIC)*, que actúa en Asia y la región del Pacífico; el *American Registry for Internet Numbers (ARIN)*, responsable por la región de América del Norte; el *Latin American and Caribbean Internet Addresses Registry (LACNIC)*, que actúa en América Latina y algunas islas del Caribe; y el *Réseaux IP Européens Network Coordination Centre (RIPE NCC)*, que sirve a Europa y los países de Asia Central.

Un Registro Nacional de Internet (*National Internet Registry - NIR*) es responsable por la asignación de direcciones IP a nivel nacional, distribuyéndolas a los Registros Locales de Internet (*Local Internet Registry - LIR*). Los LIR son generalmente ISPs que a su vez pueden ofrecer estas direcciones a los usuarios finales o a otros ISPs.

En Brasil, el NIR responsable por la distribución de direcciones IP y el registro de nombres de dominio que utilizan “.br” es el Núcleo de Informação e Coordenação do ponto br (**NIC.br**).

¿Por qué utilizar IPv6 hoy?

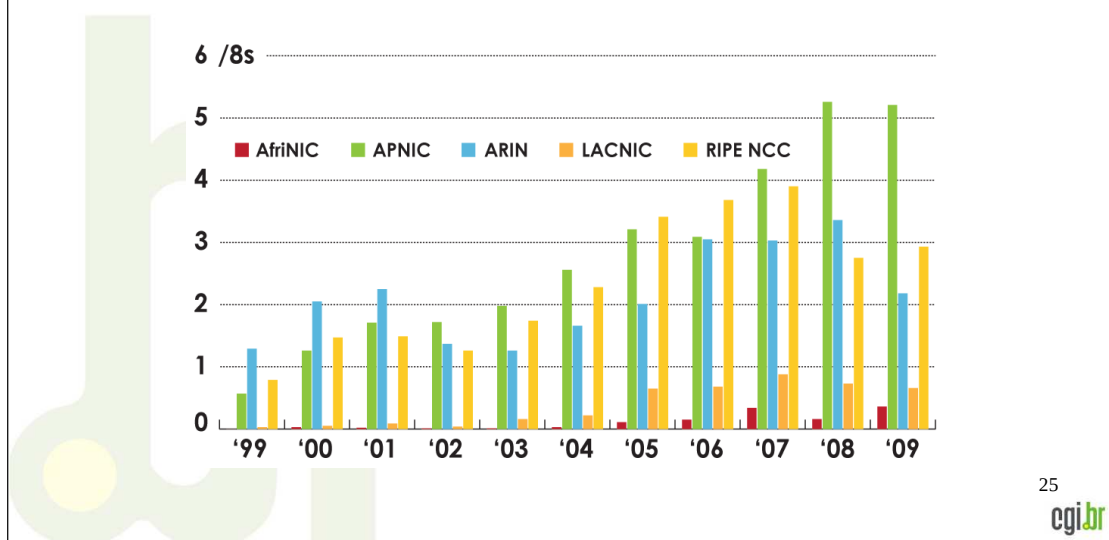
- Evolución del stock de bloques IP en la IANA.



Este gráfico muestra la evolución del stock central de bloques IP de la IANA a lo largo de los dos últimos años, junto con una proyección para los dos años siguientes.

¿Por qué utilizar IPv6 hoy?

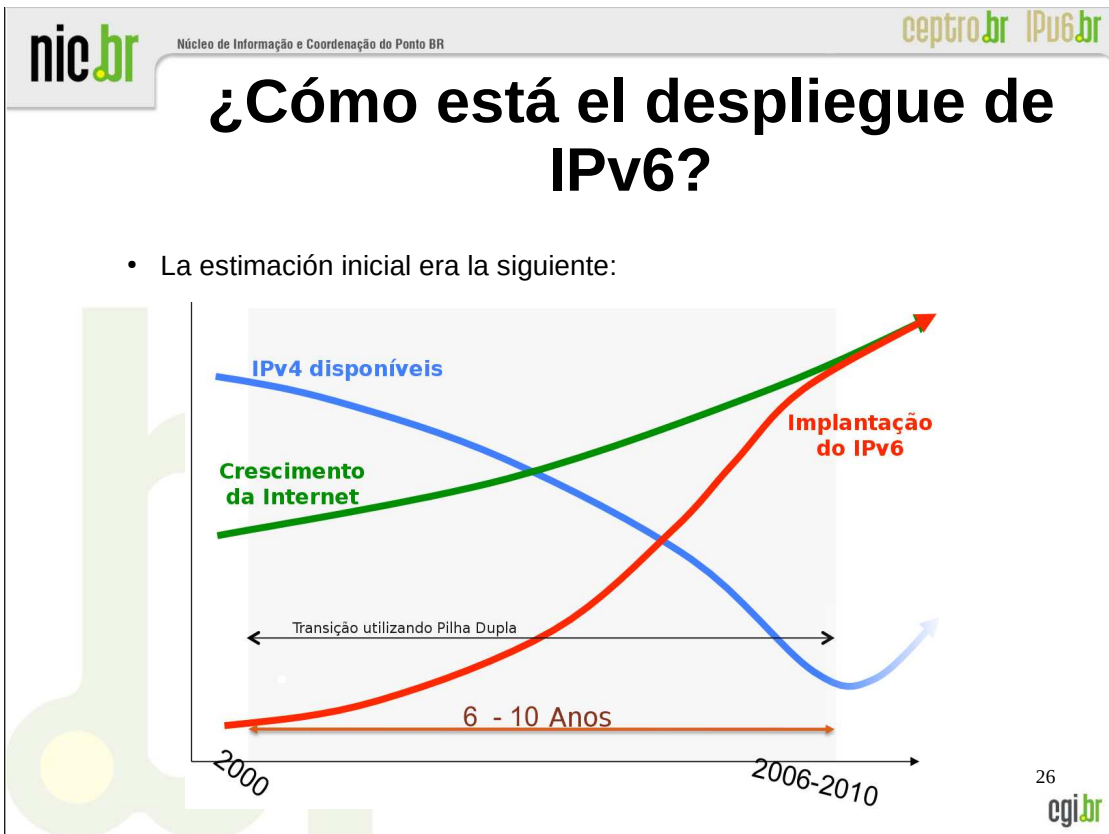
- Número de bloques (/8) IPv4 asignados anualmente por los RIR.



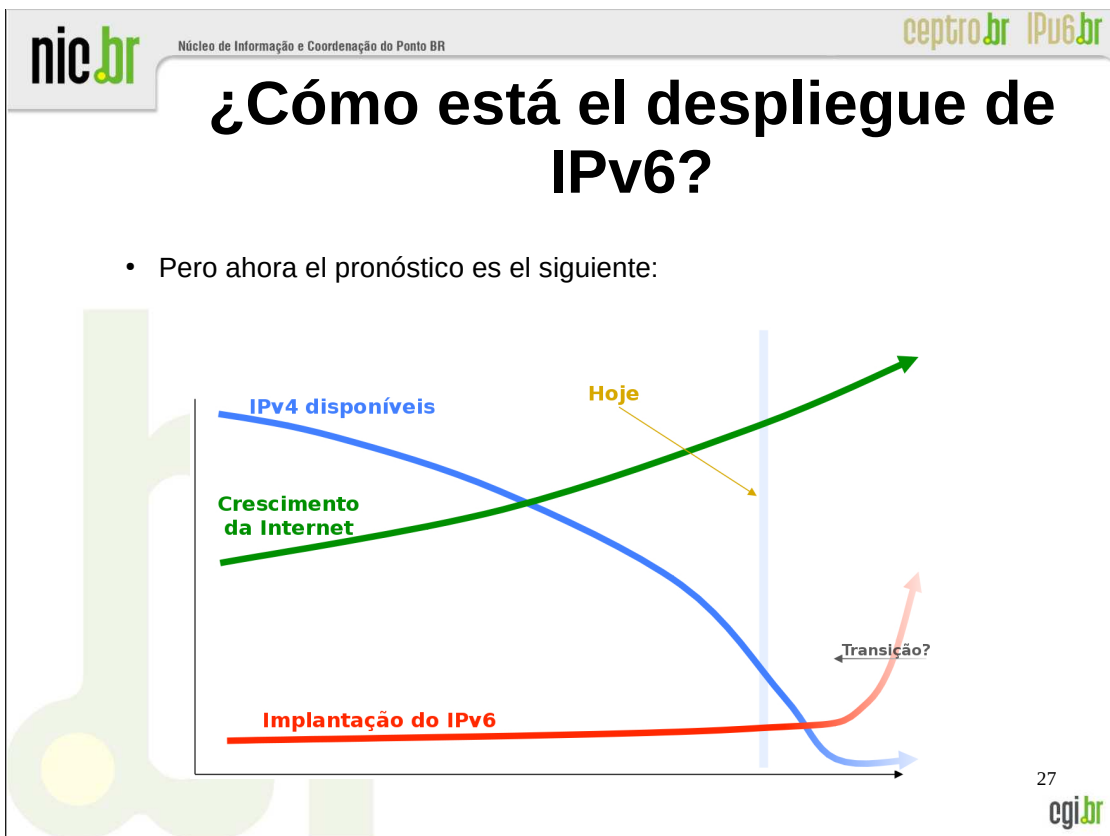
Este gráfico muestra el número de bloques /8 IPv4 asignados anualmente por los RIR.

Más información:

- <http://www.internetworldstats.com/stats.htm>
- <http://cetic.br/usuarios/ibopec/tab02-04.htm>
- <http://www.cisco.com/web/BR/barometro/barometro.html>
- <https://www.isc.org/solutions/survey>
- <http://www.nro.net/statistics>
- <http://www.abiresearch.com>



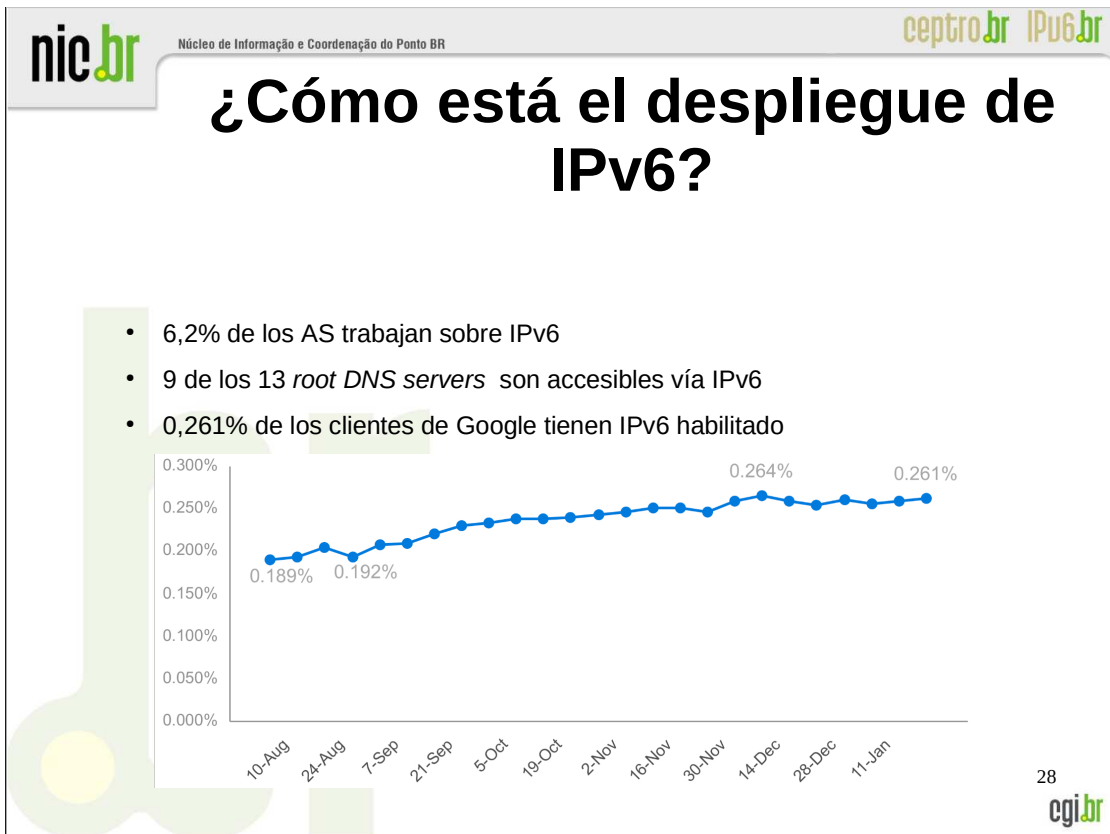
Si bien todas las cifras confirman la necesidad de contar con más direcciones IP, tema que se resuelve fácilmente con la adopción de IPv6, el despliegue de la nueva versión del protocolo IP no está ocurriendo tan rápidamente como se estimó al inicio de su desarrollo, cuando se preveía que IPv6 sería el protocolo estándar de Internet aproximadamente diez años después de su definición. En otras palabras, si eso realmente hubiese ocurrido, el objetivo de este curso probablemente sería otro.



Todavía hay muchos debates en torno al despliegue de IPv6 y algunos factores han retrasado la implementación del nuevo protocolo.

A pesar de darnos tiempo para desarrollar IPv6, las técnicas como NAT y DHCP contribuyen a demorar su adopción. A esto se suma el hecho de que IPv4 no presenta graves problemas de funcionamiento.

También debemos destacar que hasta el momento el uso de IPv6 está principalmente asociado al área académica y que, para que Internet pase a utilizar IPv6 a gran escala, es necesario que la infraestructura de los principales ISPs sea capaz de transmitir tráfico IPv6 de forma nativa. Sin embargo, su implementación en las grandes redes se ha enfrentado con dificultades, entre otras causas, debido al temor que generan los grandes cambios, los gastos debidos a la necesidad de cambiar equipos tales como routers y switches, y los gastos relacionados con el aprendizaje y la capacitación para el área técnica.



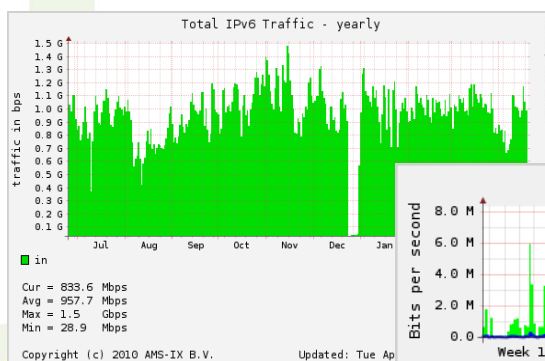
Se están realizando diferentes estudios intentando medir la cantidad de información que trafica en Internet sobre el protocolo IPv6. Análisis del número de sistemas autónomos anunciando IPv6, análisis de consultas a servidores DNS y número de páginas en Internet que utilizan IPv6 son algunos ejemplos de cómo se está intentando medir la evolución del despliegue de la versión 6 del Protocolo de Internet.

Google ha realizado una evaluación del estado actual del uso de IPv6 por parte de los usuarios comunes, recolectando información proporcionada por los navegadores de una parte de los usuarios de sus servicios. Este estudio permitió determinar que aproximadamente 0,2% de sus clientes tienen IPv6 habilitado y que la cantidad de accesos utilizando IPv6 aumentó de 0,189% en agosto de 2008 a 0,261% en enero de 2009.

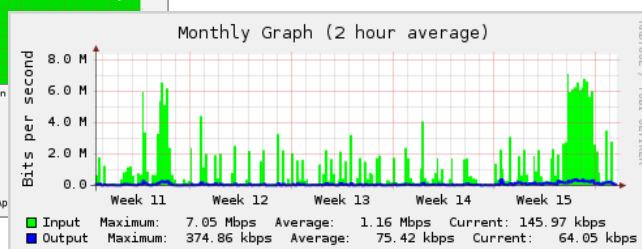
Otros datos interesantes indican que solamente el 6,2% de los sistemas autónomos trabajan sobre IPv6. También se destaca que 9 de los 13 *root DNS servers* son accesibles vía IPv6 (A, B, F, H, I, J, K, L y M).

¿Cómo está el despliegue de IPv6?

- Al menos 23% de los PTT alrededor del mundo intercambian tráfico IPv6
- En AM-IX el tráfico IPv6 intercambiado es de aproximadamente 1Gbps



- PTTMetro-SP ofrece tránsito IPv6 experimental gratuito a sus participantes

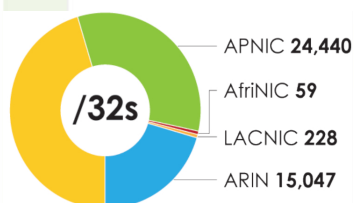


Punto fundamental de la infraestructura de Internet, el 23% de los PTT (puntos de intercambio de tráfico o, en inglés, IXP - Internet eXchange Point) alrededor del mundo intercambian tráfico IPv6, y en uno de los mayores IXP (AM-IX, Amsterdam Internet Exchange) el tráfico IPv6 intercambiado es de aproximadamente 1Gbps, que equivale a 0,3% del tráfico total.

En Brasil, desde febrero de 2010 NIC.br ofrece a los participantes de PTTMetro São Paulo el servicio de tránsito IPv6 experimental en forma gratuita. Con esta iniciativa NIC.br busca promover el uso del protocolo, reduciendo el tiempo entre la asignación de los bloques a las entidades y su uso efectivo, permitiendo la implementación y facilitando su implementación.

¿Cómo está el despliegue de IPv6?

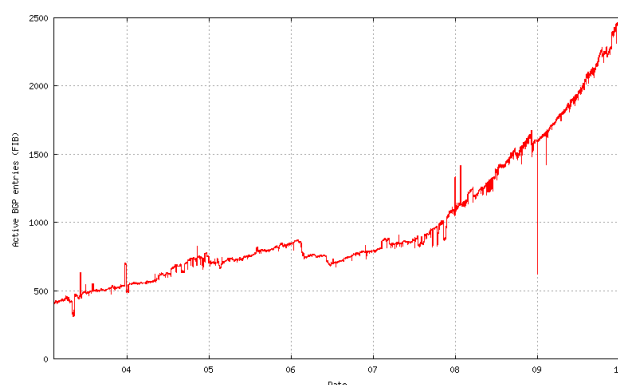
- De los ~73.000 bloques /32 ya distribuidos por los RIR, solamente un 3% se utilizan efectivamente.



RIPE NCC
33,629

Distribuciones realizadas por los RIR

Datos del 15/01/2010

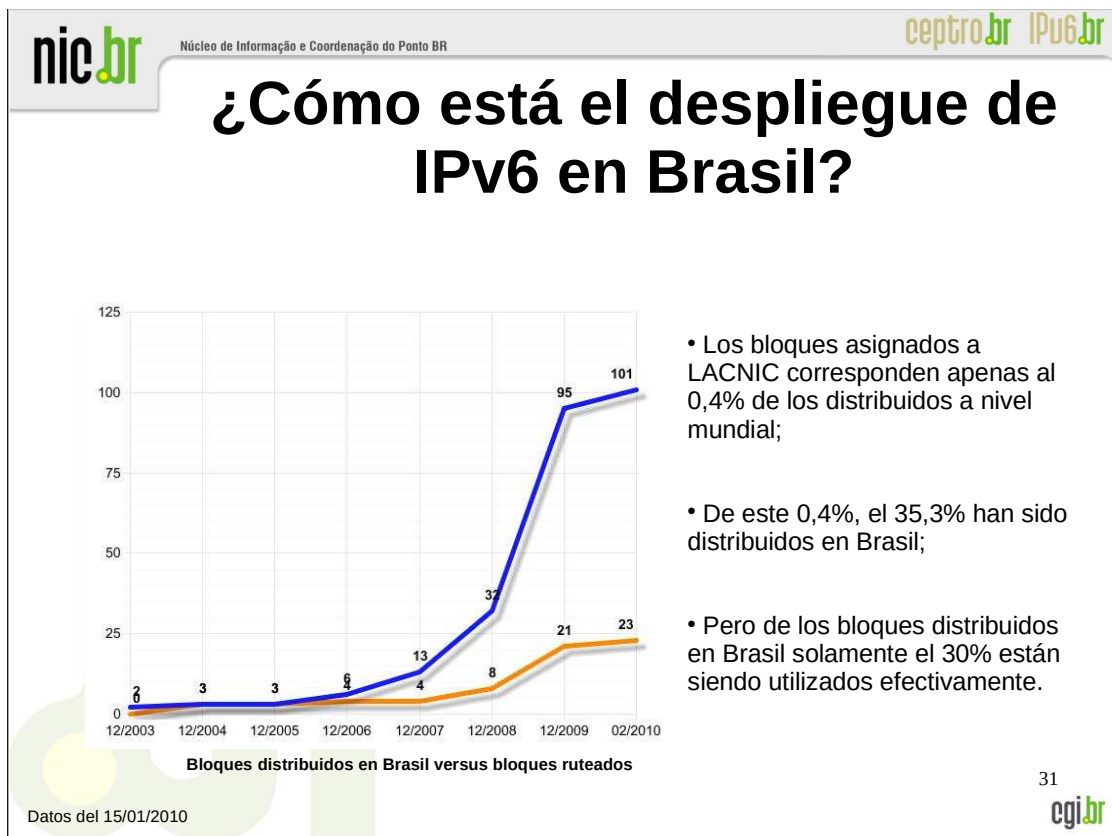


Entradas IPv6 en la tabla de enrutamiento global

30

egi.br

Hace aproximadamente diez años que los RIR están distribuyendo bloques de direcciones IPv6. Sin embargo, el hecho de que los RIR distribuyan direcciones a los Registros Nacionales o a los ISP no significa que estas direcciones estén siendo utilizadas. Cruzando los datos sobre el número de bloques /32 IPv6 ya distribuidos contra el número de rutas anunciadas en la tabla de enrutamiento se observa que apenas 3% de estos recursos están siendo utilizados efectivamente, es decir, de los 73.000 bloques ya distribuidos apenas un poco más de 3.000 están presentes en la tabla de enrutamiento global.



En Brasil, y en toda América Latina, la situación es similar.

LACNIC, el RIR que actúa en América Latina y el Caribe, ya distribuyó aproximadamente 230 bloques /32 IPv6, que corresponde a aproximadamente 0,4% del total de bloques ya distribuidos mundialmente. De ellos, el 35,3% fueron distribuidos en Brasil, pero apenas 30% están siendo utilizados efectivamente.

Más información:

- <http://www.arbornetworks.com/IPv6research>
- https://sites.google.com/site/ipv6implementors/conference2009/agenda/10_Lees_Google_IPv6_User_Measurement.pdf
- <http://www.oecd.org/dataoecd/48/51/44953210.pdf>
- <http://www.ipv6.br/IPV6/MenuIPv6Transito>
- <http://www.ams-ix.net/sflow-stats/ipv6/>
- <http://bgp.he.net/ipv6-progress-report.cgi>
- <http://portalipv6.lacnic.net/pt-br/ipv6/estat-sticas>
- <http://bgp.potaroo.net/v6/as2.0/index.html>
- <ftp://ftp.registro.br/pub/stats/delegated-ipv6-nicbr-latest>

¿Cuáles son los riesgos de no implementar IPv6?

- Aunque todavía es pequeño, el uso de IPv6 viene aumentando gradualmente;
- Pero debe avanzar aun más;
- La no implementación de IPv6:
 - Impedirá el surgimiento de nuevas redes;
 - Reducirá el proceso de inclusión digital reduciendo el número de nuevos usuarios;
 - Dificultará el surgimiento de nuevas aplicaciones;
 - Aumentará el uso técnicas como NAT.
- El costo de no implementar IPv6 puede ser mayor que el costo de implementarlo;
- Los Proveedores de Internet necesitan innovar y ofrecer nuevos servicios a sus clientes.

32

Es importante observar que, aunque el uso de IPv6 aun no es demasiado representativo, todos los datos presentados muestran que su penetración en las redes está aumentando gradualmente. Pero es necesario avanzar aun más. Postergar la implementación de IPv6 puede generar diferentes perjuicios para el desarrollo de Internet en su totalidad.

Como vimos, actualmente hay una demanda muy grande por más direcciones IP y aunque Internet pudiera seguir funcionando sin nuevas direcciones tendrá muchas dificultades para crecer. Todos los días surgen nuevas redes gracias a la expansión de las empresas y al surgimiento de nuevos negocios; numerosas iniciativas de inclusión digital han acercado nuevos usuarios a Internet; el crecimiento de las redes 3G y el uso de Internet en dispositivos electrónicos y electrodomésticos son ejemplos de nuevas aplicaciones que contribuyen al crecimiento de la red.

No implementar IPv6 probablemente impedirá el desarrollo de todas estas áreas; además, IPv6 elimina la necesidad de utilizar NAT y favorece el funcionamiento de numerosas aplicaciones. Es por ello que el costo de no utilizar el nuevo protocolo o continuar postergando su implementación será mucho mayor que el de utilizarlo.

Para los Proveedores de Internet es importante ofrecer nuevos servicios a sus clientes, principalmente porque la innovación es la clave para mantenerse delante de la competencia.

IPv6.br

La nueva generación del
Protocolo de Internet

Encabezado de IPv6

Módulo 2

34

A partir de este momento comenzaremos a estudiar las principales características de IPv6, empezando por el análisis de los cambios que ocurren en la estructura de su encabezado, presentando las diferencias entre los encabezados IPv4 e IPv6, y de qué manera esos cambios mejoran el funcionamiento del protocolo. También detallaremos el funcionamiento de los encabezados de extensión, mostrando por qué su utilización puede mejorar el desempeño de los routers.

Encabezado de IPv4

Versão (Version)	Tamanho do Cabeçalho (IHL)	Tipo de Serviço (ToS)		Tamanho Total (Total Length)	
Identificação (Identification)			Flags	Deslocamento do Fragmento (Fragment Offset)	
Tempo de Vida (TTL)		Protocolo (Protocol)		Soma de verificação do Cabeçalho (Checksum)	
Endereço de Origem (Source Address)					
Endereço de Destino (Destination Address)					
Opções + Complemento (Options + Padding)					

Un encabezado de IPv4 está formado por 12 campos fijos, que pueden o no tener opciones, por lo que su tamaño puede variar entre 20 y 60 bytes.

35

Un encabezado de IPv4 está formado por 12 campos fijos, que pueden o no tener opciones, por lo que su tamaño puede variar entre 20 y 60 bytes. Estos campos se utilizan para transmitir información sobre:

- la versión del protocolo;
- el tamaño del encabezado y los datos;
- la fragmentación;
- el tipo de datos;
- el tiempo de vida del paquete;
- el protocolo de la capa siguiente (TCP, UDP, ICMP);
- la integridad de los datos;
- el origen y el destino del paquete.

Encabezado de IPv6

- Más simple
 - 40 bytes (tamaño fijo).
 - Solo dos veces mayor que en la versión anterior.
- Más flexible
 - Extensión por medio de encabezados adicionales.
- Más eficiente
 - Minimiza el *overhead* en los encabezados.
 - Reduce el costo de procesamiento de los paquetes.

Se realizaron algunos cambios en el formato del encabezado base de IPv6 para volverlo más simple (solo ocho campos y un tamaño fijo de 40 bytes), más flexible y más eficiente, previendo su extensión por medio de encabezados adicionales que no necesitan ser procesados por todos los routers intermedios. Estos cambios permitirán que, incluso con un espacio de direccionamiento de 128 bits, cuatro veces mayor que los 32 bits de IPv4, el tamaño total del encabezado de IPv6 sea apenas dos veces mayor que el de la versión anterior.



Núcleo de Informação e Coordenação do Ponto BR



Encabezado de IPv6

Versão (Version)	Tamanho do Cabeçalho (H.L.)	Tipo de Serviço (ToS)	Tamanho Total (Total Length)	
Identificação (Identification)		Flags	Deslocamento do Fragmento (Fragment Offset)	
Tempo de Vida (TTL)	Protocolo (Protocol)	Soma de verificação do Cabeçalho (Checksum)		
Endereço de Origem (Source Address)				
Endereço de Destino (Destination Address)				
Opções + Complemento (Options + Padding)				

Versão (Version)	Classe de Tráfego (Traffic Class)	Identificador de Fluxo (Flow Label)	
Tamanho dos Dados (Payload Length)		Próximo Cabeçalho (Next Header)	Limite de Encaminhamento (Hop Limit)
Endereço de Origem (Source Address)			
Endereço de Destino (Destination Address)			

- Se eliminaron seis campos del encabezado de IPv4.

37


Entre estos cambios se destaca la eliminación de seis campos del encabezado de IPv4 debido a que sus funciones ya no son necesarias o bien son implementadas por los encabezados de extensión.

Ahora las opciones adicionales forman parte de los encabezados de extensión de IPv6. Por lo tanto se pudieron eliminar los campos Opciones y Complementos.

El campo Tamaño del Encabezado también se eliminó, ya que el tamaño del encabezado de IPv6 es fijo.

Los campos Identificación, *Flags* y Desplazamiento del Fragmento se eliminaron, ya que los datos referentes a la fragmentación ahora se indican en un encabezado de extensión apropiado.

Para aumentar la velocidad de procesamiento de los routers se eliminó el campo Suma de Verificación, ya que este cálculo es realizado por los protocolos de las capas superiores.

Núcleo de Informação e Coordenação do Ponto BR

Encabezado de IPv6

Versão (Version)	Tamanho do Cabeçalho (H/L)	1 Tipo de Serviço (ToS)	2 Tamanho Total (Total Length)	
Identificação (Identification)		Flags	Deslocamento do Fragmento (Fragment Offset)	
4 Tempo de Vida (TTL)	3 Protocolo (Protocol)	Soma de verificação do Cabeçalho (Checksum)		
Endereço de Origem (Source Address)				
Endereço de Destino (Destination Address)				
Opções + Complemento (Options + Padding)				

Versão (Version)	1 Classe de Tráfego (Traffic Class)	Identificador de Fluxo (Flow Label)	
2 Tamanho dos Dados (Payload Length)		3 Próximo Cabeçalho (Next Header)	4 Limite de Encaminhamento (Hop Limit)
Endereço de Origem (Source Address)			
Endereço de Destino (Destination Address)			

- Se eliminaron seis campos del encabezado de IPv4.
- Los nombres de cuatro campos fueron modificados, al igual que sus ubicaciones.

38

Otro cambio fufe el cambio de nombre y de ubicación de otros cuatro campos.

IPv4	IPv6
Tipo de Servicio	→ Clase de Tráfico
Tamaño Total	→ Tamaño de los Datos
Tiempo de Vida (TTL)	→ Límite de Encaminamiento
Protocolo	→ Siguiete Encabezado

Estos cambios de posición se definieron para facilitar el procesamiento de estos datos por parte de los routers.

Encabezado de IPv6

Versão (Version)	Tamanho do Cabeçalho (H-L)	Tipo de Serviço (ToS)		Tamanho Total (Total Length)		Versão (Version)	Classe de Tráfego (Traffic Class)	Identificador de Fluxo (Flow Label)	
Identificação (Identification)				Flags	Deslocamento do Fragmento (Fragment Offset)		Endereço de Origem (Source Address)		
Tempo de Vida (TTL)		Protocolo (Protocol)		Soma de verificação do Cabeçalho (Checksum)					
Endereço de Origem (Source Address)									
Endereço de Destino (Destination Address)									
Opções + Complemento (Options + Padding)									

Versão (Version)	Tamanho dos Dados (Payload Length)		Próximo Cabeçalho (Next Header)	Limite de Encaminhamento (Hop Limit)
Endereço de Destino (Destination Address)				

- Se eliminaron seis campos del encabezado de IPv4.
- Los nombres de cuatro campos fueron modificados, al igual que sus ubicaciones.
- Se agregó el campo Identificador de Flujo.

También se agregó un nuevo campo, el Identificador de Flujo, agregándole al protocolo IP otro mecanismo de soporte para QoS. En los próximos módulos de este curso presentaremos más detalles sobre este campo y cómo el protocolo IPv6 trata el tema de QoS.

Encabezado de IPv6

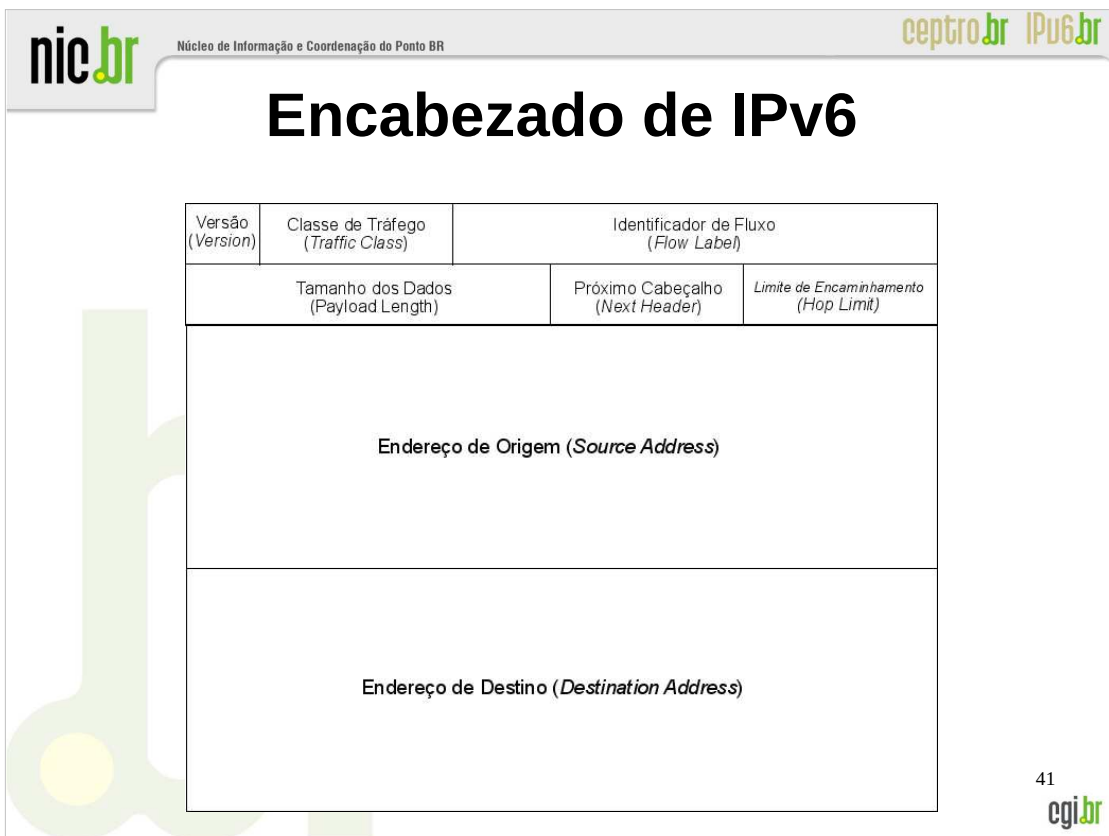
Versão (Version)	Tamanho do Cabeçalho (H.L)	Tipo de Serviço (ToS)	Tamanho Total (Total Length)		Versão (Version)	Classe de Tráfego (Traffic Class)	Identificador de Fluxo (Flow Label)	
Identificação (Identification)			Flags	Deslocamento do Fragmento (Fragment Offset)	Tamanho dos Dados (Payload Length)		Próximo Cabeçalho (Next Header)	Limite de Encaminhamento (Hop Limit)
Tempo de Vida (TTL)	Protocolo (Protocol)	Soma de verificação do Cabeçalho (Checksum)						
Endereço de Origem (Source Address)								
Endereço de Destino (Destination Address)								
Opções + Complemento (Options + Padding)								

Endereço de Origem (Source Address)								
Endereço de Destino (Destination Address)								

- Se eliminaron seis campos del encabezado de IPv4.
- Los nombres de cuatro campos fueron modificados, al igual que sus ubicaciones.
- Se agregó el campo Identificador de Flujo.
- Se mantuvieron tres campos.

40

Los campos Versión, Dirección de Origen y Dirección de destino se mantuvieron, modificando solamente el tamaño del espacio reservado para el direccionamiento que pasó a tener 128 bits.

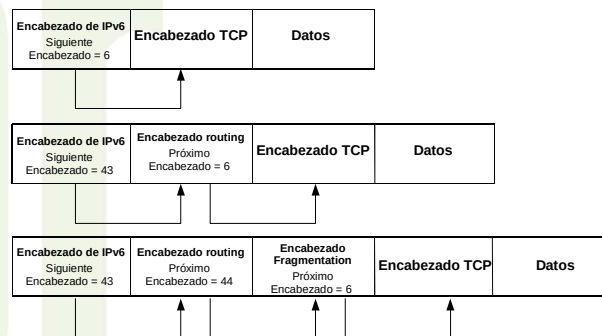


Vamos a conocer un poco acerca de cada campo del encabezado base de IPv6:

- **Versión** (4 bits) - Identifica la versión del protocolo IP utilizado. En el caso de IPv6 el valor de este campo es 6.
- **Clase de Tráfico** (8 bits) - Identifica y diferencia los paquetes por clases de servicios o prioridad. Este continúa ofreciendo las mismas funcionalidades y definiciones del campo Tipo de Servicio de IPv4.
- **Identificador de Flujo** (20 bits) - Identifica y diferencia paquetes del mismo flujo en la capa de red. Este campo permite que el router identifique el tipo de flujo de cada paquete, sin necesidad de verificar su aplicación.
- **Tamaño de los Datos** (16 bits) - Indica el tamaño, en bytes, solamente de los datos enviados junto con el encabezado de IPv6. Reemplaza al campo Tamaño Total usado en IPv4, que indica el tamaño del encabezado más el tamaño de los datos transmitidos. En el cálculo del tamaño también se incluyen los encabezados de extensión.
- **Siguiente Encabezado** (8 bits) - Identifica el encabezado que sigue al encabezado de IPv6. El nombre de este campo fue modificado (en IPv4 se llamaba Protocolo) para reflejar la nueva organización de los paquetes IPv6, ya que ahora este campo no solo contiene valores referentes a otros protocolos sino que también indica los valores de los encabezados de extensión.
- **Límite de Encaminamiento** (8 bits) - Indica el número máximo de routers que el paquete IPv6 puede pasar antes de ser descartado; se decrementa en cada salto. Estandarizó el modo en que se utilizaba el campo Tiempo de Vida (TTL) de IPv4, a pesar de la definición original del campo TTL, diciendo que éste debe indicar, en segundos, el tiempo que el paquete demorará en ser descartado en caso de no llegar a su destino.
- **Dirección de Origen** (128 bits) - Indica la dirección de origen del paquete.
- **Dirección de Destino** (128 bits) - Indica la dirección de destino del paquete.

Encabezados de extensión

- En IPv6 las opciones se tratan por medio de los encabezados de extensión.
- Éstos se encuentran en el encabezado base y el encabezado de la capa de transporte.
- Estos encabezados no tienen ni cantidad ni tamaño fijo.



42

A diferencia de lo que ocurre en IPv4, donde todos los datos opcionales se incluyen en el encabezado base, en IPv6 estos datos se incluyen a través de encabezados de extensión. Estos encabezados se encuentran entre el encabezado base y el encabezado de la capa inmediatamente superior, y no tienen una cantidad ni un tamaño fijo. Si en un mismo paquete existen múltiples encabezados de extensión, éstos serán agregados en serie formando una “cadena de encabezados”.

Las especificaciones de IPv6 definen seis encabezados de extensión: *Hop-by-Hop Options*, *Destination Options*, *Routing*, *Fragmentation*, *Authentication Header* y *Encapsulating Security Payload*.

En IPv6, el uso de encabezados de extensión busca aumentar la velocidad de procesamiento en los routers, ya que el único encabezado de extensión procesado en cada router es el *Hop-by-Hop*; los demás solo son tratados por el nodo identificado en el campo Dirección de Destino del encabezado base. Además, se pueden definir y utilizar nuevos encabezados de extensión sin tener que modificar el encabezado base.

Encabezados de extensión

Hop-by-Hop Options

- Identificado por el valor 0 en el campo Siguiente Encabezado.
- Transporta datos que deben ser procesados por todos los nodos a lo largo del camino del paquete.

Siguiente Encabezado	Tam. encab. de extensión	
		Opciones

43

Identificado por el valor 0 en el campo Siguiente Encabezado, el encabezado de extensión *Hop-by-Hop* se debe colocar inmediatamente después del encabezado base de IPv6. Los datos transmitidos por el mismo deben ser examinados por todos los nodos intermedios a lo largo del camino del paquete hasta que llega a su destino. En su ausencia, el router sabe que no necesita procesar ningún dato adicional y puede encaminar el paquete hacia el destino final inmediatamente.

Las siguientes son las definiciones de cada campo del encabezado:

- **Siguiente Encabezado** (1 byte): Identifica el tipo de encabezado que sigue al *Hop-by-Hop*.
- **Tamaño del Encabezado** (1 byte): Indica el tamaño del encabezado *Hop-by-Hop* en unidades de 8 bytes, excluyendo los ocho primeros.
- **Opciones**: Contiene una o más opciones y su tamaño es variable. En este campo, el primer bite contiene información acerca de cómo estas opciones deben ser tratadas en caso que el nodo que está procesando la información no la reconozca. El valor de los primeros dos bits especifica las acciones a tomar:
 - 00: ignorar y continuar el procesamiento.
 - 01: descartar el paquete.
 - 10: descartar el paquete y enviar un mensaje ICMP *Parameter Problem* a la dirección de origen del paquete.
 - 11: descartar el paquete y enviar un mensaje ICMP *Parameter Problem* a la dirección de origen del paquete, solamente si el destino no es una dirección *multicast*.

El tercer bit de este campo especifica si la información opcional puede cambiar de ruta (valor 01) o no (valor 00).

Hasta el momento existen dos tipos definidos para el encabezado *Hop-by-Hop: Router Alert* y *Jumbogram*.

- ***Router Alert***: Se utiliza para informar a los intermediarios que el mensaje a ser encaminado exige tratamiento especial. Esta opción es utilizada por los protocolos MLD (*Multicast Listener Discovery*) y RSVP (*Resource Reservation Protocol*).
- ***Jumbogram***: Se utiliza para informar que el tamaño del paquete IPv6 es mayor que 64KB.

Más información:

- RFC 2711 - *IPv6 Router Alert Option*

Encabezados de extensión

Destination Options

- Identificado por el valor 60 en el campo Siguiente Encabezado.
- Transporta datos que deben ser procesados por el nodo de destino del paquete.

Siguiente Encabezado	Tam. encab. de extensión	
		Opciones

45

Identificado por el valor 60 en el campo Siguiente Encabezado, el encabezado de extensión *Destination Options* transporta datos que deben ser procesados por el nodo de destino del paquete, indicado en el campo Dirección de Destino del encabezado base. La definición de sus campos es igual a la de los encabezado *Hop-by-Hop*.

Este encabezado se utiliza en el soporte para movilidad en IPv6 a través de la opción *Home Address*, la cual contiene la Dirección de Origen del Nodo Móvil cuando está en tránsito.

Encabezados de extensión

Routing

- Identificado por el valor 43 en el campo Siguiente Encabezado.
- Inicialmente desarrollado para listar uno o más nodos intermedios a ser visitados por el paquete antes de llegar a su destino.
- Actualmente utilizado como parte del mecanismo de soporte para movilidad en IPv6.

Siguiente Encabezado	Tam. encab. de extensión	Tipo de Routing	Salto restantes
Reservado			
Dirección de origen			

46

Identificado por el valor 43 en el campo Siguiente Encabezado, el encabezado de extensión *Routing* fue inicialmente desarrollado para listar uno o más nodos intermedios a ser visitados por el paquete antes de llegar a su destino, de manera similar a las opciones *Loose Source* y *Record Route* de IPv4. Esta función, realizada por el encabezado *Routing Type 0*, se volvió obsoleta con la RFC5095 debido a problemas de seguridad.

Se definió un nuevo encabezado *Routing, Type 2*, para ser utilizado como parte del mecanismo de soporte para movilidad en IPv6, transportando la Dirección de Origen del Nodo Móvil en paquetes enviados por el Nodo Correspondiente.

Las siguientes son las definiciones de cada campo del encabezado:

- **Siguiente Encabezado** (1 byte): Identifica el tipo de encabezado que sigue al encabezado *Routing*.
- **Tamaño del Encabezado** (1 byte): Indica el tamaño del encabezado *Routing* en unidades de 8 bytes, excluyendo los ocho primeros.
- **Routing Type** (1 byte): Identifica el tipo de encabezado *Routing*. Actualmente solo está definido el *Type 2*.
- **Salto restantes**: Definido para ser utilizado con *Routing Type 0*, indica el número de saltos a ser visitados antes que el paquete llegue a su destino final.
- **Dirección de Origen**: Transporta la Dirección de Origen de un Nodo Móvil.

Más información:

- RFC 3775 - *Mobility Support in IPv6 - 6.4. Type 2 Routing Header*
- RFC 5095 - *Deprecation of Type 0 Routing Headers in IPv6*

Encabezados de extensión

Fragmentation

- Identificado por el valor 44 en el campo Siguiente Encabezado.
- Transporta información sobre los fragmentos de los paquetes IPv6.

Siguiente Encabezado	Reservado	Desplazamiento del Fragmento	Res	M
Identificación				

47

egi.br

Identificado por el valor 44 en el campo Siguiente Encabezado, el encabezado de extensión *Fragmentation* se utiliza cuando el paquete IPv6 a ser enviado es mayor que la *Path MTU*.

Las siguientes son las definiciones de cada campo del encabezado:

- **Siguiente Encabezado** (1 byte): Identifica el tipo de encabezado que sigue al encabezado *Fragmentation*.
- **Desplazamiento del Fragmento** (13 bits): Indica, en unidades de ocho bytes, la posición de los datos transportados por el fragmento actual respecto del inicio del paquete original.
- **Flag M** (1 bit): Si está marcado con el valor 1, indica que hay más fragmentos. Si está marcado con el valor 0, indica que es el fragmento final.
- **Identificación** (4 bytes): Valor único generado por el nodo de origen para identificar el paquete original. Se utiliza para detectar los fragmentos de un mismo paquete.

El proceso de fragmentación de paquetes de IPv6 se describirá más detalladamente en los próximos módulos.

Encabezados de extensión

Authentication Header

- Identificado por el valor 51 en el campo Siguiente Encabezado.
- Utilizado por IPSec para proveer autenticación y garantía de integridad a los paquetes IPv6.

Encapsulating Security Payload

- Identificado por el valor 52 en el campo Siguiente Encabezado.
- También utilizado por IPSec, garantiza la integridad y confidencialidad de los paquetes.

Los encabezados de extensión *Authentication Header* y *Encapsulating Security Payload*, indicados respectivamente por los valores 51 y 52 en el campo Siguiente Encabezado, forman parte del encabezado IPSec.

Aunque las funcionalidades de IPSec son idénticas tanto en IPv4 como en IPv6, su utilización con IPv6 es facilitada por el hecho de que sus principales elementos forman parte integral de la nueva versión del protocolo IP. También hay otros aspectos que facilitan esta utilización, como el hecho de no utilizar NAT con IPv6, tema que se discutirá en los próximos módulos junto con los detalles de los encabezados de extensión AH y ESP.

Encabezados de extensión

- Cuando hay más de un encabezado de extensión se recomienda que aparezcan en el siguiente orden:
 - *Hop-by-Hop Options*
 - *Routing*
 - *Fragmentation*
 - *Authentication Header*
 - *Encapsulating Security Payload*
 - *Destination Options*
- Si el campo Dirección de Destino tiene una dirección *multicast*, los encabezados de extensión serán examinados por todos los nodos del grupo.
- El encabezado de extensión *Mobility* puede ser utilizado por quienes cuentan con soporte para movilidad en IPv6 .

49

Observemos algunos aspectos de los encabezados de extensión.

En primer lugar es importante destacar que, para evitar que los nodos existentes a lo largo del camino del paquete tengan que recorrer toda la cadena de encabezados de extensión para saber cuáles datos tratar, estos encabezados se deben enviar respetando un orden determinado. En general, los encabezados importantes para todos los nodos involucrados en el enrutamiento se deben colocar en primer lugar, los encabezados que solo son importantes para el destinatario final se colocan al final de la cadena. La ventaja de esta secuencia es que el nodo puede detener el procesamiento de los encabezados apenas encuentre algún encabezado de extensión dedicado al destino final, teniendo la certeza de que no hay más encabezados importantes. De este modo se puede mejorar significativamente el procesamiento de paquetes, ya que en muchos casos el procesamiento del encabezado base será suficiente para encaminar el paquete. Por lo tanto, la secuencia a seguir sería la siguiente:

- *Hop-by-Hop Options*
- *Routing*
- *Fragmentation*
- *Authentication Header*
- *Encapsulating Security Payload*
- *Destination Options*

También vale la pena observar que si un paquete fue enviado a una dirección *multicast*, los encabezados de extensión serán examinados por todos los nodos del grupo.

En cuanto a la flexibilidad que ofrecen los encabezados de extensión, vale la pena destacar el desarrollo del encabezado *Mobility*, utilizado por los nodos que poseen soporte para movilidad en IPv6.

IPv6.br

La nueva generación del
Protocolo de Internet

Direccionamiento IPv6

Módulo 3

51

La principal característica y mayor justificación para el desarrollo del protocolo IPv6 es el aumento del espacio de direccionamiento. Por este motivo es importante conocer las diferencias entre las direcciones IPv4 e IPv6, saber reconocer la sintaxis de las direcciones IPv6 y conocer los tipos de direcciones IPv6 que existen y sus principales características.

Direccionamiento

- Una dirección IPv4 está formada por 32 bits.

$$2^{32} = 4.294.967.296$$

- Una dirección IPv6 está formada por 128 bits.

$$2^{128} = \mathbf{340.282.366.920.938.463.463.374.607.431.768.211.456}$$

~ $5,6 \times 10^{28}$ direcciones IP por cada ser humano.

~ $7,9 \times 10^{28}$ de direcciones más que en IPv4.

En IPv4, el campo del encabezado reservado para direccionamiento tiene 32 bits. Este tamaño permite un máximo de 4.294.967.296 (2^{32}) direcciones diferentes. En la época de su desarrollo, esta cantidad se consideraba suficiente para identificar todas las computadoras en la red y soportar el surgimiento de nuevas subredes. Sin embargo, con el rápido crecimiento de la Internet surgió el problema de la escasez de las direcciones IPv4, que motivó la creación de una nueva generación del protocolo IP.

IPv6 tiene un espacio de direccionamiento de 128 bits, lo que permite obtener 340.282.366.920.938.463.463.374.607.431.768.211.456 direcciones (2^{128}). Este valor representa aproximadamente $7,9 \times 10^{28}$ de direcciones más que IPv4, y más de $5,6 \times 10^{28}$ de direcciones por cada ser humano en la Tierra si consideramos una población estimada en 6 mil millones de habitantes.

Direccionamiento

La representación de las direcciones IPv6 divide la dirección en ocho grupos de 16 bits, separados mediante “:”, escritos con dígitos hexadecimales.

2001:0DB8:AD1F:25E2:CADE:CAFE:F0CA:84C1

2 bits

En la representación de una dirección IPv6 está permitido:

- Utilizar caracteres en mayúscula o minúscula;
- Omitir los ceros a la izquierda; y
- Representar los ceros continuos mediante “::”.

Ejemplo:

2001:0DB8:0000:0000:130F:0000:0000:140B

2001:db8:0:0:130f::140b

Formato no válido: **2001:db8::130f::140b** (genera ambigüedad)

53

Los 32 bits de las direcciones IPv4 se dividen en cuatro grupos de 8 bits cada uno, separados por “.”, escritos con dígitos decimales. Por ejemplo: **192.168.0.10**.

La representación de las direcciones IPv6 divide la dirección en ocho grupos de 16 bits, separados mediante “:”, escritos con dígitos hexadecimales (0-F). Por ejemplo:

2001:0DB8:AD1F:25E2:CADE:CAFE:F0CA:84C1

En la representación de las direcciones IPv6 está permitido utilizar tanto caracteres en mayúscula como en minúscula.

Además, se pueden aplicar reglas de abreviatura para facilitar la escritura de algunas direcciones muy extensas. Se permite omitir los ceros a la izquierda de cada bloque de 16 bits y también reemplazar una larga secuencia de ceros por “::”. Por ejemplo, la dirección **2001:0DB8:0000:0000:130F:0000:0000:140B** se puede escribir como **2001:DB8:0:0:130F::140B** o **2001:DB8::130F:0:0:140B**. En este ejemplo se puede observar que la abreviatura del grupo de ceros solo se puede realizar una vez, caso contrario podría provocar ambigüedad en la representación de la dirección. Si la dirección anterior se escribiera como **2001:DB8::130F::140B**, no sería posible determinar si corresponde a **2001:DB8:0:0:130F:0:0:140B**, a **2001:DB8:0:0:0:130F:0:140B** o **2001:DB8:0:130F:0:0:0:140B**.

Esta abreviatura también se puede realizar al final o al inicio de la dirección, como ocurre en **2001:DB8:0:54:0:0:0:0** que se puede escribir como **2001:DB8:0:54::**.

Direccionamiento

- Representación de los prefijos
 - Como CIDR (IPv4)
 - “dirección-IPv6/tamaño del prefijo”
 - Ejemplo:

Prefijo **2001:db8:3003:2::/64**
Prefijo global **2001:db8::/32**
ID de la subred **3003:2**
- URL
 - [http://\[2001:12ff:0:4::22\]/index.html](http://[2001:12ff:0:4::22]/index.html)
 - [http://\[2001:12ff:0:4::22\]:8080](http://[2001:12ff:0:4::22]:8080)

54

Otra representación importante es la de los prefijos de red. En las direcciones IPv6 continúa escribiéndose del mismo modo que en IPv4, utilizando la notación CIDR. Esta notación se representa con la forma “dirección-IPv6/tamaño del prefijo”, donde “tamaño del prefijo” es un valor decimal que especifica la cantidad de bits contiguos a la izquierda de la dirección que comprenden el prefijo. El ejemplo de prefijo de subred presentado a continuación que, de los 128 bits de la dirección, 64 bits se utilizan para identificar la subred.

Prefijo **2001:db8:3003:2::/64**
Prefijo global **2001:db8::/32**
ID de la subred **3003:2**

Esta representación también permite agregar las direcciones en forma jerárquica, identificando la topología de la red a través de parámetros tales como ubicación geográfica, proveedor de acceso, identificación de la red, división de la subred, etc. Esto permite disminuir el tamaño de la tabla de enrutamiento y agilizar el encaminamiento de los paquetes.

En cuanto a la representación de las direcciones IPv6 en las URL (*Uniform Resource Locators*), éstas ahora se representan entre corchetes. Esto evitará ambigüedades en caso que sea necesario indicar el número de un puerto junto con la URL. Veamos los siguientes ejemplos:

[http://\[2001:12ff:0:4::22\]/index.html](http://[2001:12ff:0:4::22]/index.html)
[http://\[2001:12ff:0:4::22\]:8080](http://[2001:12ff:0:4::22]:8080)

Direccionamiento

En IPv6 se han definido tres tipos de direcciones:

- **Unicast** → Identificación Individual
- **Anycast** → Identificación Selectiva
- **Multicast** → Identificación en Grupo

No existen más las direcciones **Broadcast**.

55

cgi.br

En IPv6 se han definido tres tipos de direcciones:

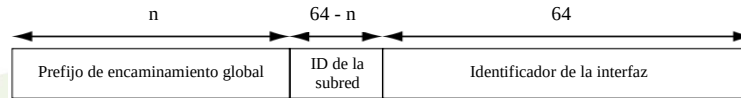
- **Unicast** – Este tipo de dirección identifica una única interfaz, de modo que un paquete enviado a una dirección *unicast* se entrega a una única interfaz;
- **Anycast** – Identifica un conjunto de interfaces. Un paquete enviado a una dirección *anycast* se entrega a la interfaz perteneciente a este conjunto más próxima al origen (de acuerdo con la distancia medida por los protocolos de encaminamiento). Las direcciones *anycast* se utilizan en comunicaciones de "uno" a "uno-de-muchos".
- **Multicast** – También identifica un conjunto de interfaces, pero un paquete enviado a una dirección *multicast* se entrega a todas las interfaces asociadas a esa dirección. Las direcciones *multicast* se utilizan en comunicaciones de "uno" a "muchos".

A diferencia de lo que ocurre en IPv4, en IPv6 no existe la dirección *broadcast*, responsable de direccionar un paquete a todos los nodos de un mismo dominio. En IPv6 esta función ha sido asignada a determinados tipos de direcciones *multicast*.

Direccionamiento

Unicast

- *Global Unicast*



- **2000::/3**

- Globalmente ruteable (similar a las direcciones IPv4 públicas);
- 13% del total de direcciones posibles;
- $2^{(45)} = 35.184.372.088.832$ redes /48 diferentes.

56

Las direcciones *unicast* se utilizan para comunicaciones entre dos nodos, por ejemplo, teléfonos VoIPv6, computadoras en una red privada, etc., y su estructura fue definida para permitir agregaciones con prefijos de tamaño flexible, similar a CIDR en IPv4.

Existen varios tipos de direcciones *unicast* IPv6: *Global Unicast*; *Unique-Local*; y *Link-Local* por ejemplo. También existen algunos tipos de direcciones para usos especiales, entre ellas las direcciones IPv4 mapeadas en direcciones IPv6, las direcciones de *loopback* y las direcciones no especificadas.

- ***Global Unicast*** - Equivalente a las direcciones IPv4 públicas, las direcciones *global unicast* son globalmente ruteables y accesibles en la Internet IPv6. Está formada por tres partes: el prefijo de encaminamiento global, utilizado para identificar el tamaño del bloque atribuido a una red; la identificación de la subred, utilizada para identificar un enlace en una red; y la identificación de la interfaz, que debe identificar de forma única una interfaz dentro de un enlace.

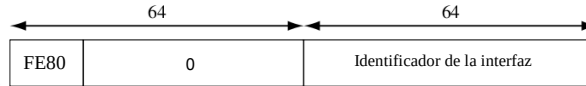
Su estructura fue proyectada para utilizar los 64 bits más hacia la izquierda para identificar la red y los 64 bits más hacia la derecha para identificar la interfaz. Por lo tanto, excepto en ciertos casos específicos, todas las subredes en IPv6 tienen el mismo tamaño de prefijo, 64 bits (/64), lo que permite tener $2^{64} = 18.446.744.073.709.551.616$ dispositivos por subred.

Actualmente para la atribución de direcciones está reservado el rango **2000::/3** (001), que corresponde a las direcciones de **2000::** a **3fff:ffff:ffff:ffff:ffff:ffff:ffff:ffff**. Esto representa 13% del total de direcciones posibles con IPv6, lo que permite crear $2^{(64-3)} = 2.305.843.009.213.693.952$ ($2,3 \times 10^{18}$) subredes (/64) diferentes o $2^{(48-3)} = 35.184.372.088.832$ ($3,5 \times 10^{13}$) redes /48.

Direccionamiento

Unicast

- *Link local*



- **FE80::/64**

- Solo se debe utilizar localmente;
- Atribuido automáticamente (autoconfiguración *stateless*);

57

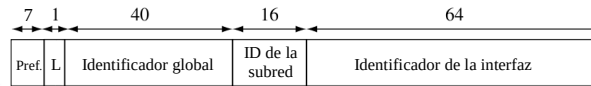
egi.br

- **Link Local** – Pudiendo utilizarse solo en el enlace específico en el cual la interfaz está conectada, la dirección *link local* es atribuida automáticamente usando el prefijo **FE80::/64**. Los 64 bits reservados para la identificación de la interfaz se configuran utilizando el formato IEEE EUI-64. Vale la pena destacar que los routers no deben encaminar paquetes cuyo origen o destino sea una dirección *link-local* hacia otros enlaces;

Direccionamiento

Unicast

- *Unique local*



- **FC00::/7**

- Prefijo globalmente único (con alta probabilidad de ser único);
- Se utiliza solo en las comunicaciones dentro de un enlace o entre un conjunto limitado de enlaces;
- No se espera que sea ruteado en Internet.

58

- **Unique Local Address (ULA)** – Dirección con grandes probabilidades de ser globalmente única, utilizada solamente para comunicaciones locales, generalmente dentro de un mismo enlace o conjunto de enlaces. Una dirección ULA no debe ser ruteable en la Internet global.

Una dirección ULA, creada utilizando un ID global distribuido pseudo-aleatoriamente, está formada por las siguientes partes:

- **Prefijo: FC00::/7.**
- **Flag Local (L):** si el valor es 1 (**FD**) el prefijo es atribuido localmente. Si el valor es 0 (**FC**), el prefijo debe ser atribuido por una organización central (aun por determinar).
- **Identificador global:** identificador de 40 bits usado para crear un prefijo globalmente único.
- **Identificador de la interfaz:** identificador de la interfaz de 64 bits.

De este modo, la estructura de una dirección ULA es **FDUU:UUUU:UUUU:<ID de la subred>:<ID de la interfaz>** donde **U** son los bits del identificador único, generado aleatoriamente por un algoritmo específico.

Su utilización permite que cualquier enlace posea un prefijo /48 privado y globalmente único. Por lo tanto, en el caso que se interconecten dos redes (por ejemplo de dos empresas diferentes), es probable que no haya conflicto de direcciones ni necesidad de reenumerar la interfaz. Además, la dirección ULA es independiente del proveedor, pudiendo ser utilizado en la comunicación dentro del enlace aunque no haya conexión a Internet. Otra ventaja es que su prefijo se puede bloquear fácilmente, y si accidentalmente una dirección ULA se anuncia fuera del enlace ya sea a través de un router o vía DNS, no habrá conflicto con otras direcciones.

Direccionamiento

Unicast

- Identificador de la Interfaz (IID)
 - Deben ser únicos dentro del mismo prefijo de subred.
 - El mismo IID se puede usar en múltiples interfaces de un único nodo ya que se refieren a subredes diferentes.
 - Normalmente se utiliza un IID de 64 bits, que se puede obtener:
 - Manualmente
 - Autoconfiguración *stateless*
 - DHCPv6 (*stateful*)
 - A partir de una clave pública (CGA)
 - El IID puede ser temporario y generado aleatoriamente.
 - Normalmente se basa en la dirección MAC (Formato EUI-64).

59

Los identificadores de interfaz (IID), utilizados para distinguir las interfaces dentro de un enlace, deben ser únicos dentro del mismo prefijo de subred. El mismo IID se puede usar en múltiples interfaces de un único nodo, pero éstas deben estar asociadas a subredes diferentes.

Normalmente se utiliza un IID de 64 bits, el cual se puede obtener de diferentes maneras: Se puede configurar manualmente, a partir del mecanismo de autoconfiguración *stateless* de IPv6, a partir de servidores DHCPv6 (*stateful*), o se pueden formar a partir de una clave pública (CGA). Estos métodos se describirán detalladamente en este curso.

Aunque pueden ser generados aleatoriamente y de forma temporaria, se recomienda construir el IID en base a la dirección MAC de la interfaz, en el formato EUI-64.

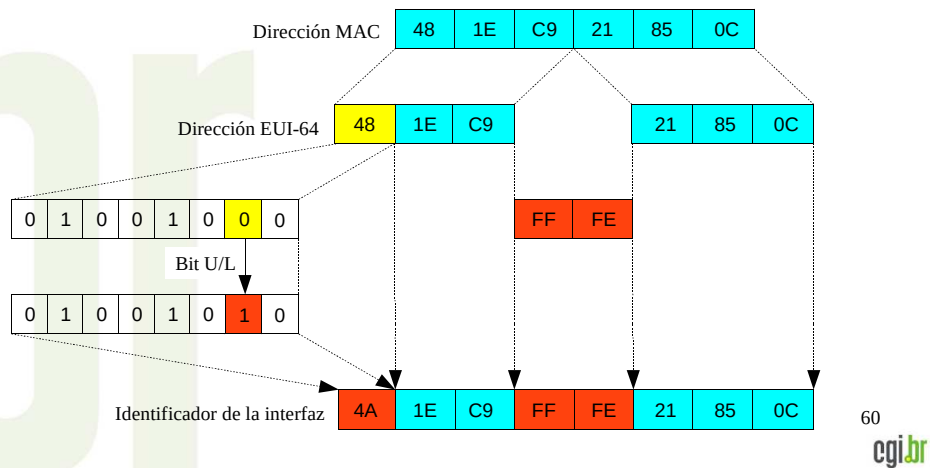
Más información:

- RFC 3986 - *Uniform Resource Identifier (URI): Generic Syntax*
- RFC 4291 - *IP Version 6 Addressing Architecture*
- RFC 4193 - *Unique Local IPv6 Unicast Addresses*
- RFC 5156 - *Special-Use IPv6 Addresses*
- RFC 3587 - *IPv6 Global Unicast Address Format*
- *Internet Protocol Version 6 Address Space* - <http://www.iana.org/assignments/ipv6-address-space>

Direccionamiento

Unicast

- EUI-64



Un IID basado en el formato EUI-64 se genera de la siguiente manera:

- Si la interfaz tiene una dirección MAC de 64 bits (estándar EUI-64), solo debe complementar el séptimo bit más a la izquierda (llamado bit U/L – Universal/Local) de la dirección MAC, es decir, si es 1 se debe cambiar a 0 y si es 0 se debe cambiar a 1. Si la interfaz utiliza una dirección MAC de 48 bits (estándar IEEE 802), primero se agregan los dígitos hexadecimales FF-FE entre el tercero y el cuarto byte de la dirección MAC (para transformar al estándar EUI-64), y luego se complementa el bit U/L. Por ejemplo:
- Si la dirección MAC de la interfaz es:
 - 48-1E-C9-21-85-0C
- se agregan los dígitos FF-FE en el medio de la dirección:
 - 48-1E-C9-FF-FE-21-85-0C
- se complementa el bit U/L:
 - 48 = 01001000
 - 01001000 → 01001010
 - 01001010 = 4A
- IID = 4A-1E-C9-FF-FE-21-85-0C

Una dirección *link local* atribuida a esta interfaz sería **FE80::4A1E:C9FF:FE21:850C**.

Más información:

- Guidelines for 64-bit Global Identifier (EUI-64) Registration Authority* - <http://standards.ieee.org/regauth/oui/tutorials/EUI64.html>

Direccionamiento

Unicast

- Direcciones especiales
 - Localhost - **::1/128 (0:0:0:0:0:0:0:1)**
 - No especificada - **::/128 (0:0:0:0:0:0:0:0)**
 - mapeada IPv4 - **::FFFF:wxyz**
- Rangos especiales
 - 6to4 - **2002::/16**
 - Documentación - **2001:db8::/32**
 - Teredo - **2001:0000::/32**
- Obsoletas
 - Site local - **FEC0::/10**
 - IPv4-compatible - **::wxyz**
 - 6Bone - **3FFE::/16** (red de prueba desactivada el 06/06/06)

61

Existen algunas direcciones IPv6 especiales utilizadas para fines específicos:

- **Dirección no especificada (*Unspecified*):** representada mediante la dirección **0:0:0:0:0:0:0:0** o **::0** (equivalente a la dirección IPv4 *unspecified* **0.0.0.0**). Esta nunca debe ser atribuida a ningún nodo, ya que solamente indica la ausencia de una dirección. Se la puede utilizar, por ejemplo, en el campo Dirección de Origen de un paquete IPv6 enviado por un *host* durante el proceso de inicialización, antes que éste tenga determinada su dirección exclusiva. La dirección *unspecified* no debe ser utilizada como dirección de destino de paquetes IPv6;
- **Dirección *Loopback*:** representada mediante la dirección *unicast* **0:0:0:0:0:0:0:1** o **::1** (equivalente a la dirección IPv4 *loopback* **127.0.0.1**). Esta dirección se utiliza para referenciar la propia máquina y es muy utilizada para realizar pruebas internas. Este tipo de dirección no debe atribuirse a ninguna interfaz física ni ser utilizada como dirección de origen en paquetes IPv6 enviados a otros nodos. Además, un paquete IPv6 con una dirección *loopback* como destino no puede ser enviado por un router IPv6, y si un paquete recibido en una interfaz tiene una dirección *loopback* como destino, éste debe ser descartado;
- **Direcciones mapeadas IPv4:** representada mediante **0:0:0:0:0:FFFF:wxyz** o **::FFFF:wxyz**, se utiliza para mapear una dirección IPv4 en una dirección IPv6 de 128 bits, donde **wxyz** representa los 32 bits de la dirección IPv4, utilizando dígitos decimales. Se aplica en técnicas de transición para que se comuniquen nodos IPv6 e IPv4. Por ejemplo **::FFFF:192.168.100.1**.

También se han reservado algunos rangos de direcciones para usos específicos:

- **2002::/16**: prefijo utilizado en el mecanismo de transición 6to4;
- **2001:0000::/32**: prefijo utilizado en el mecanismo de transición TEREDO;
- **2001:db8::/32**: prefijo utilizado para representar direcciones IPv6 en textos y documentaciones.

Otras direcciones que se utilizaron en los inicios del desarrollo de IPv6 se han vuelto obsoletas y ya no deben ser utilizadas:

- **FEC0::/10**: prefijo utilizado por las direcciones tipo *site local*, desarrolladas para ser utilizadas dentro de una red específica sin necesidad de un prefijo global, equivalente a las direcciones privadas en IPv4. Su utilización fue reemplazada por las direcciones ULA;
- **::wxyz**: utilizado para representar la dirección IPv4-compatible. Su función es la misma que la de la dirección mapeada IPv4, y se ha tornado obsoleto debido a su caída en desuso;
- **3FFE::/16**: prefijo utilizado para representar las direcciones de la red de prueba 6Bone. Creada para ayudar en la implementación de IPv6, esta red se desactivó el 6 de junio de 2006 (06/06/06).

Más información:

- RFC 3849 - *IPv6 Address Prefix Reserved for Documentation*
- RFC 3879 - *Deprecating Site Local Addresses*

Direccionamiento

Anycast

- Identifica un grupo de interfaces
 - Entrega el paquete solo a la interfaz más cercana al origen.
- Atribuidas a partir de direcciones *unicast* (son iguales desde el punto de vista sintáctico).
- Posibles usos:
 - Descubrir servicios en la red (DNS, *proxy* HTTP, etc.);
 - Balanceo de carga;
 - Localizar routers que proveen acceso a una determinada subred;
 - Utilizado en redes con soporte para movilidad IPv6 para localizar los Agentes de Origen...
- *Subnet-Router*

63

cgi.br

Una dirección IPv6 *anycast* se utiliza para identificar un grupo de interfaces, aunque con la propiedad de que un paquete enviado a una dirección *anycast* es encaminado solamente a la interfaz del grupo más próxima al origen del paquete.

Las direcciones *anycast* se atribuyen a partir del rango de direcciones *unicast* y no hay diferencias sintácticas entre las mismas. Por lo tanto, una dirección *unicast* atribuida a más de una interfaz se transforma en una dirección *anycast*, debiéndose en este caso configurar explícitamente los nodos para que sepan que les ha sido atribuida una dirección *anycast*. Por otra parte, esta dirección se debe configurar en los routers como una entrada independiente (prefijo /128 – *host route*).

Este esquema de direccionamiento se puede utilizar para descubrir servicios en la red, como por ejemplo servidores DNS y *proxies* HTTP, garantizando la redundancia de estos servicios. También se puede utilizar para realizar balanceo de carga en situaciones donde múltiples *hosts* o routers proveen el mismo servicio, para localizar routers que provean acceso a una determinada subred o para localizar los Agentes de Origen en redes con soporte para movilidad IPv6.

Todos los routers deben soportar la dirección *anycast Subnet-Router*. Este tipo de direcciones están formadas por el prefijo de la subred y el IID completado con ceros (por ejemplo: **2001:db8:cafe:dad0::/64**). Un paquete enviado a la dirección *Subnet-Router* será entregada al router más próximo al origen dentro de la misma subred.

También se definió una dirección *anycast* para ser utilizada en el soporte para movilidad IPv6. Este tipo de direcciones están formadas por el prefijo de la subred seguido por el IID **ffff:ffff:ffff:fffe** (por ejemplo: **2001:db8::ffff:ffff:ffff:fffe**). Son utilizadas por el Nodo Móvil cuando éste necesita localizar un Agente Origen en su Red Original.

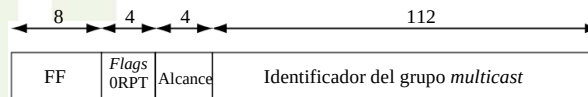
Más información:

- *Internet Protocol Version 6 Anycast Addresses* - <http://www.iana.org/assignments/ipv6-anycast-addresses>

Direccionamiento

Multicast

- Identifica un grupo de interfaces.
- El soporte para *multicast* es obligatorio en todos los nodos IPv6.
- La dirección *multicast* deriva del bloque **FF00::/8**.
- El prefijo **FF** es seguido por cuatro bits utilizados como *flags* y otros cuatro bits que definen el alcance de la dirección *multicast*. Los 112 bits restantes se utilizan para identificar el grupo *multicast*.



64

cgi.br

Las direcciones *multicast* se utilizan para identificar grupos de interfaces, ya que cada interfaz puede pertenecer a más de un grupo. Los paquetes enviados a esas direcciones se entregan a todas las interfaces que componen el grupo.

En IPv4, el soporte *multicast* es opcional, ya que fue introducido solamente como una extensión al protocolo. Sin embargo, en IPv6 se requiere que todos los nodos soporten *multicast*, ya que muchas funcionalidades de la nueva versión del protocolo IP utilizan este tipo de direcciones.

Su funcionamiento es similar al de *broadcast*, dado que un único paquete es enviado a varios *hosts*, diferenciándose solo por el hecho de que en *broadcast* el paquete se envía a todos los *hosts* de la red sin excepción, mientras que en *multicast* solamente un grupo de *hosts* reciben el paquete. De este modo la posibilidad de transportar solo una copia de los datos a todos los elementos del grupo, a partir de un árbol de distribución puede reducir la utilización de recurso en una red, además de optimizar la entrega de datos a los *hosts* receptores. Aplicaciones tales como las videoconferencias, distribución de video bajo demanda, actualizaciones de *software* y juegos *on-line*, son ejemplos de servicios que vienen ganando notoriedad y pueden aprovechar las ventajas de *multicast*.

Las direcciones *multicast* no deben ser utilizadas como dirección de origen de un paquete. Estas direcciones derivan del bloque **FF00::/8**, donde el prefijo **FF**, que identifica una dirección *multicast*, es precedido por cuatro bits, que representan cuatro *flags*, y un valor de cuatro bits que define el alcance del grupo *multicast*. Los 112 bits restantes se utilizan para identificar el grupo *multicast*.

Direccionamiento

Multicast

Flags

Flag	Valor (binario)	Descripción
Primer bit	0	Marcado como 0 (Reservado para uso futuro)
R	1	Dirección de un Punto de Encuentro (<i>Rendezvous Point</i>)
R	0	No representa una dirección de Punto de Encuentro
P	1	Dirección <i>multicast</i> basada en el prefijo de red
P	0	Dirección <i>multicast</i> no basada en el prefijo de red
T	1	Dirección <i>multicast</i> temporaria (no distribuida por IANA)
T	0	Dirección <i>multicast</i> permanente (distribuida por IANA)

Alcance

Valor (4 bits hex)	Descripción
1	Interfaz
2	Enlace
3	Subred
4	Admin
5	Site
8	Organización
E	Global
(0, F)	Reservados
(6, 7, 9, A, B, C, D)	No distribuidos

65

egi.br

Las *flags* se definen de la siguiente manera:

- El primer bit más a la izquierda está reservado y debe ser marcado con 0;
- **Flag R:** Si el valor es 1 indica que la dirección *multicast* “lleva” la dirección de un Punto de Encuentro (*Rendezvous Point*). Si el valor es 0 indica que no lleva una dirección de Punto de Encuentro integrada;
- **Flag P:** Si el valor es 1 indica que la dirección *multicast* está basada en un prefijo de red. Si el valor es 0 indica que la dirección no está basada en un prefijo de red;
- **Flag T:** Si el valor es 0 indica que la dirección *multicast* es permanente, es decir que ha sido atribuida por la IANA. Si el valor es 1 indica que la dirección *multicast* no es permanente, es decir que ha sido atribuida dinámicamente.

Los cuatro bits que representan el alcance de la dirección *multicast* se utilizan para delimitar el área que abarca un grupo *multicast*. Los valores atribuidos a este campo son los siguientes:

- 1 - solo abarca la interfaz local;
- 2 - abarca los nodos de un enlace;
- 3 - abarca los nodos de una subred;
- 4 - abarca la menor área que se puede configurar manualmente;
- 5 - abarca los nodos de un site;
- 8 - abarca varios sites de una misma organización;
- E - abarca toda la Internet;
- 0, F - reservados;
- 6, 7, 9, A, B, C, D - no están distribuidos.

Así, un router conectado al *backbone* de la Internet no encaminará paquetes con un alcance menor que 14 (E en hexa), por ejemplo. En IPv4, el alcance de un grupo *multicast* se especifica a través del campo TTL del encabezado.

Direccionamiento

Multicast

Dirección	Alcance	Descripción
FF01::1	Interfaz	Todas las interfaces (<i>all-nodes</i>)
FF01::2	Interfaz	Todos los routers (<i>all-routers</i>)
FF02::1	Enlace	Todos los nodos (<i>all-nodes</i>)
FF02::2	Enlace	Todos los routers (<i>all-routers</i>)
FF02::5	Enlace	Routers OSFP
FF02::6	Enlace	Routers OSPF designados
FF02::9	Enlace	Routers RIP
FF02::D	Enlace	Routers PIM
FF02::1:2	Enlace	Agentes DHCP
FF02::1:FFXX:XXXX	Enlace	<i>Solicited-node</i>
FF05::2	Site	Todos los routers (<i>all-routers</i>)
FF05::1:3	Site	Servidores DHCP en un site
FF05::1:4	Site	Agentes DHCP en un site
FF0X::101	Variado	NTP (<i>Network Time Protocol</i>)

66

La siguiente lista muestra algunas direcciones *multicast* permanentes:

Dirección	Alcance	Descripción
FF01::1	Interfaz	Todas las interfaces en un nodo (<i>all-nodes</i>)
FF01::2	Interfaz	Todos los routers en un nodo (<i>all-routers</i>)
FF02::1	Enlace	Todos los nodos del enlace (<i>all-nodes</i>)
FF02::2	Enlace	Todos los routers del enlace (<i>all-routers</i>)
FF02::5	Enlace	Routers OSFP
FF02::6	Enlace	Routers OSPF designados
FF02::9	Enlace	Routers RIP
FF02::D	Enlace	Routers PIM
FF02::1:2	Enlace	Agentes DHCP
FF02::1:FFXX:XXXX	Enlace	<i>Solicited-node</i>
FF05::2	Site	Todos los routers en un site
FF05::1:3	Site	Servidores DHCP en un site
FF05::1:4	Site	Agentes DHCP en un site
FF0X::101	Variado	NTP (<i>Network Time Protocol</i>)

Direccionamiento

Multicast

- Dirección *Solicited-Node*
 - Todos los nodos deben formar parte de este grupo;
 - Se forma agregando el prefijo **FF02::1:FF00:0000/104** a los 24 bits más a la derecha del IID;
 - Utilizado por el protocolo de Descubrimiento de Vecinos (*Neighbor Discovery*).

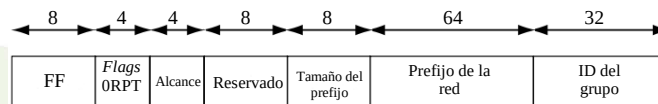
La dirección *multicast solicited-node* identifica un grupo *multicast* del cual todos los nodos pasan a formar parte una vez que les es asignada una dirección *unicast* o *anycast*. Una dirección *solicited-node* se forma agregando al prefijo **FF02::1:FF00:0000/104** los 24 bits más a la derecha del identificador de la interfaz, y para cada dirección *unicast* o *anycast* del nodo existe una dirección *multicast solicited-node* correspondiente.

En las redes IPv6, la dirección *solicited-node* es utilizada por el protocolo de Descubrimiento de Vecinos para resolver la dirección MAC de una interfaz. Para ello se envía un mensaje *Neighbor Solicitation* a la dirección *solicited-node*. De este modo solo las interfaces registradas en este grupo examinan el paquete. En una red IPv4, para determinar la dirección MAC de una interfaz se envía un mensaje *ARP Request* a la dirección *broadcast* de la capa de enlace, de modo que todas las interfaces del enlace examinan el mensaje.

Direccionamiento

Multicast

- Dirección *multicast* derivada de un prefijo *unicast*



- *Flag* P = 1
- *Flag* T = 1
- Prefijo **FF30::/12**
- Ejemplo:

prefijo de red	=	2001:DB8::/32
dirección	=	FF3E:20:2001:DB8:0:0:CADE:CAFE

68

Con el propósito de reducir el número de protocolos necesarios para la distribución de direcciones *multicast*, se definió un formato extendido de dirección *multicast*, que permite distribuir direcciones en base a prefijos *unicast* y direcciones SSM (*source-specific multicast*).

En las direcciones basadas en el prefijo de red, el *flag* P se marca con el valor 1. En este caso el uso del campo alcance no influye, aunque el alcance de esta dirección *multicast* no debe exceder el alcance del prefijo *unicast* “cargado” junto al mismo. Los 8 bits posteriores al campo alcance están reservados y deben ser marcados con ceros. En la secuencia hay 8 bits que especifican el tamaño del prefijo de red indicado en los 64 bits siguientes. Si el prefijo de red es menor que 64 bits, los bits no utilizados en el campo tamaño de prefijo deben ser marcados con ceros. El campo identificador del grupo utiliza los 32 bits restantes. Obsérvese que en una dirección donde el *flag* P está marcado con el valor 1, el *flag* T también se debe marcar con el valor 1, ya que ésta no representa una dirección definida por la IANA.

Direccionamiento

Multicast

- Direcciones *Multicast* SSM
 - Prefijo: **FF3X::/32**
 - Formato de la dirección: **FF3X::/96**
 - Tamaño del prefijo = 0
 - Prefijo = 0
 - Ejemplo: **FF3X::CADE:CAFE/96**
donde **X** es el alcance y **CADE:CAFE** es el identificador del grupo.

En el modelo tradicional de *multicast*, llamado *any-source multicast* (ASN), un participante de un grupo *multicast* no controla de qué fuente desea recibir los datos. Con SSM una interfaz se puede registrar en un grupo *multicast* y especificar las fuentes de datos. SSM se puede implementar utilizando el protocolo MLDv2 (*Multicast Listener Discovery versión 2*).

Para una dirección SSM, los *flags* P y T se marcan con valor 1. Los campos tamaño de prefijo y prefijo de red se marcan con ceros, llegando al prefijo **FF3X::/32**, donde **X** es el valor del alcance. El campo Dirección de Origen del encabezado de IPv6 identifica al dueño de la dirección *multicast*. Toda dirección SSM tiene el formato **FF3X::/96**.

Los métodos de gestión de grupos *multicast* se discutirán en el próximo módulo del curso.

Direccionamiento

- Al igual que en IPv4, las direcciones IPv6 se atribuyen a las interfaces físicas y no a los nodos.
- Con IPv6 es posible atribuir una única interfaz a múltiples direcciones, independientemente de su tipo.
- Así un nodo se puede identificar a través de cualquier dirección de sus interfaces.
 - Loopback **::1**
 - Link Local **FE80:....**
 - Unique local **FD07:...**
 - Global **2001:....**
- La RFC 3484 determina el algoritmo para seleccionar las direcciones de origen y destino.

70

También es importante destacar algunas características relacionadas con la dirección que proporciona la nueva arquitectura del protocolo IPv6. Al igual que en IPv4, las direcciones IPv6 se atribuyen a las interfaces físicas, no a los nodos, de modo que cada interfaz necesita al menos una dirección *unicast*. Sin embargo, se puede atribuir a una única interfaz múltiples direcciones IPv6, independientemente del tipo (*unicast*, *multicast* o *anycast*) o sub-tipo (*loopback*, *link local*, *6to4*, etc.). Así un nodo se puede identificar a través de cualquier dirección de sus interfaces y por lo tanto se hace necesario elegir entre sus múltiples direcciones cuáles se utilizarán como direcciones de origen y destino al establecer una conexión.

Para resolver este tema se definieron dos algoritmos, uno para seleccionar la dirección de origen y otro para seleccionar la de destino. Estos algoritmos, que deben ser implementados por todos los nodos IPv6, especifican el comportamiento por defecto de dichos nodos, pero no anulan las decisiones tomadas por las aplicaciones o protocolos de la capa superior.

Entre las reglas más importantes se destacan las siguientes:

- Pares de direcciones del mismo alcance o tipo tienen preferencia;
- El menor alcance para la dirección de destino tiene preferencia (se utiliza el menor alcance posible);
- Las direcciones cuyo tiempo de vida no ha expirado tienen preferencia sobre las direcciones con tiempo de vida expirado;
- No se pueden utilizar las direcciones de las técnicas de transición (ISATAP, 6to4, etc.) si hay una dirección IPv6 nativa disponible;
- Si todos los criterios son similares, los pares de direcciones con el mayor prefijo común tendrán preferencia;
- Para las direcciones de origen, las direcciones globales tendrán preferencia sobre las direcciones temporarias;
- En un Nodo Móvil, la Dirección de Origen tiene preferencia sobre una Dirección Remota.

Estas reglas deben ser utilizadas cuando no hay ninguna otra especificación. Las especificaciones también permiten configurar políticas que puedan reemplazar estas preferencias estándares por combinaciones de direcciones de origen y destino.

Más información:

- RFC 2375 - *IPv6 Multicast Address Assignments*
- RFC 3306 - *Unicast-Prefix-based IPv6 Multicast*
- RFC 3307 - *Allocation Guidelines for IPv6 Multicast Addresses*
- RFC 3484 - *Default Address Selection for Internet Protocol version 6 (IPv6)*
- RFC 3569 - *An Overview of Source-Specific Multicast (SSM)*
- RFC 3956 - *Embedding the Rendezvous Point (RP) Address in an IPv6 Multicast Address*
- RFC 4007 - *IPv6 Scoped Address Architecture*
- RFC 4489 - *A Method for Generating Link-Scoped IPv6 Multicast Addresses*
- *Internet Protocol Version 6 Multicast Addresses* - <http://www.iana.org/assignments/ipv6-multicast-addresses/>

Políticas de distribución y asignación

- Cada RIR recibe de la IANA un bloque /12
- El bloque 2800::/12 corresponde al espacio reservado para LACNIC – NIC.br trabaja con un /16 que forma parte de este /12
- La distribución mínima para los ISP es un bloque /32
- Se pueden realizar distribuciones mayores si se justifica la utilización
- **¡ATENCIÓN!** A diferencia de lo que ocurre en IPv4, en IPv6 la utilización se mide considerando el número de bloques de direcciones asignados a usuarios finales, no el número de direcciones asignadas a usuarios finales.

72

En la jerarquía de las políticas de atribución, distribución y asignación de direcciones, cada RIR recibe de la IANA un bloque /12 IPv6.

El bloque **2800::/12** corresponde al espacio reservado para que LACNIC realice distribuciones en América Latina. Por su parte, NIC.br trabaja con un /16 que forma parte de este /12.

La mínima distribución para un ISP es un bloque /32, aunque se pueden realizar distribuciones mayor si se justifica la utilización. Un aspecto importante a destacar es que, a diferencia de lo que ocurre en IPv4, en IPv6 la utilización se mide considerando el número de bloques de direcciones asignados a usuarios finales, no el número de direcciones asignadas a usuarios finales.

Enfoque: *one size fits all*

- Recomendaciones para la asignación de direcciones (RFC3177):
 - En general se recomienda el uso de redes /48 para todos los tipos de usuarios, sin importar si son residenciales o empresas pequeñas o grandes;
 - Las empresas muy grandes pueden recibir un /47, prefijos un poco menores o múltiplos de un /48;
 - Se recomienda el uso de redes /64 cuando exista la certeza de que se requiere una única subred, por ejemplo para usuarios 3G;
 - Se puede utilizar una red /128 cuando exista la certeza absoluta de que solamente se conectará una interfaz.

73

egi.br

En relación con la distribución y asignación de direcciones a usuarios finales, la RFC 3177 recomienda seguir un enfoque conocido como *one size fits all*, el cual tiene las siguientes características:

- En general se recomienda el uso de redes /48 para todos los tipos de usuarios, sin importar si son residenciales o empresas pequeñas o grandes;
- Las empresas muy grandes pueden recibir un /47, prefijos un poco menores o múltiplos de un /48;
- Se recomienda el uso de redes /64 cuando exista la certeza de que se requiere una única subred, por ejemplo para usuarios 3G;
- Se puede utilizar una red /128 cuando exista la certeza absoluta de que solamente se conectará una interfaz. Por ejemplo: conexiones PPoE.

Enfoque: *one size fits all*

- Facilita la renumeración de la red en caso de cambio de proveedor (cambio de prefijo);
- Permite ampliar la red sin necesidad de solicitar más direcciones al proveedor;
- Facilita el mapeo entre la dirección global y la dirección *Unique Local* (ULA **fc00:xyzw:klmn::/48**);
- Ya hay redes que utilizan prefijos /48 6to4;
- Permite mantener reglas únicas para zonas inversas de diferentes prefijos;
- Facilita la administración;
- Hay quienes creen que desperdicia demasiadas direcciones y que podría generar problemas dentro de algunas décadas.

74

El enfoque *one size fits all* tiene algunas ventajas:

- Facilita la renumeración de la red en caso de cambio de proveedor (cambio de prefijo);
- Permite ampliar la red sin necesidad de solicitar más direcciones al proveedor;
- Facilita el mapeo entre la dirección global y la dirección *Unique Local* (ULA **fc00:xyzw:klmn::/48**);
- Ya hay redes que utilizan prefijos /48 6to4;
- Permite mantener reglas únicas para zonas inversas de diferentes prefijos;
- Facilita la administración;
- Hay quienes creen que desperdicia demasiadas direcciones y que podría generar problemas dentro de algunas décadas.

Enfoque conservador

- Si usamos “*one size fits all*...”
 - un /32 permite 65.536 /48.
- No delegar bloques /48 a todos, asignando un bloque /56 a los usuarios residenciales y SOHOs.
- Reduce el consumo total de direcciones de 6 a 7 bits.

75

cgi.br

Un enfoque más conservador, opuesto a *one size fits all*, recomienda no delegar bloques /48 a todos los tipos de usuarios, asignando bloques /56 a los usuarios residenciales y SOHOs. Esto reduce el consumo total de direcciones de 6 a 7 bits.

Además, un /32 permite “apenas” 65.536 /48, que en el caso de los grandes proveedores no sería suficiente para satisfacer toda su demanda.

¿Qué están haciendo los RIR y los ISP?

- LACNIC y AFRINIC
 - Evalúan la solicitud de bloques adicionales por parte de los ISP en base a la cantidad de bloques /48 asignados.
 - *Threshold* → HD-Ratio = 0.94.
- APNIC, ARIN y RIPE
 - Evalúan la solicitud de bloques adicionales por parte de los ISP en base a la cantidad de bloques /56 asignados.
 - *Threshold* → HD-Ratio = 0.94.

$$HD = \frac{\log(\text{número de objetos distribuidos})}{\log(\text{número de objetos distribuibles})}$$

76

Los RIR aplican dos políticas diferentes en cuanto a las recomendaciones de uso para los ISP y el criterio para la distribución de bloques de direcciones adicionales.

Los RIR LACNIC y AFRINIC siguen la recomendación *one size fits all*, y sugieren que los proveedores de sus regiones también los sigan. También evalúan las solicitudes de bloques adicionales por parte de los ISP de acuerdo con ese enfoque, es decir basándose en la cantidad de bloques /48 asignados por los ISP.

En cambio los RIR APNIC, ARIN y RIPE siguen un enfoque más conservador, utilizando la cantidad de bloques /56 asignados por los proveedores como base para evaluar las solicitudes de bloques adicionales.

En todos los casos la medida utilizada para la evaluación es el HD-Ratio (*Host-Density ratio*). El HD-Ratio es un modo de medir el uso del espacio de direccionamiento, ya que su valor está relacionado con el porcentaje de uso. Para calcular el HD-Ratio se utiliza la siguiente fórmula:

$$HD = \frac{\log(\text{número de objetos distribuidos})}{\log(\text{número de objetos distribuibles})}$$

Todos los RIR utilizan como valor *Threshold* (límite) HD-Ratio = 0,94, solo que LACNIC y AFRINIC lo aplican a la utilización de bloques /48 mientras que APNIC, ARIN y RIPE a la utilización de bloques /56.

¿Qué están haciendo los RIR y los ISP?

Bloque	No. /48	Threshold (HD=0,94)	% de utilización
/32	65.536	33.689	51,41%
/31	131.072	64.634	49,31%
/30	262.144	124.002	47,30%
/29	524.288	237.901	45,38%
/28	1.048.576	456.419	43,53%
/27	2.097.152	875.653	41,75%
/26	4.194.304	1.679.965	40,05%
/25	8.388.608	3.223.061	38,42%
/24	16.777.216	6.183.533	36,86%
/23	33.554.432	11.863.283	35,36%
/22	67.108.864	22.760.044	33,92%
/21	134.217.728	43.665.787	32,53%
/20	268.435.456	83.774.045	31,21%

77

Esta tabla presenta el porcentaje de utilización de bloques /48 en base al cálculo de HD-Ratio igual a 0,94.

¿Qué están haciendo los RIR y los ISP?

Bloque	No. /56	Threshold (HD=0,94)	% de utilización
/32	16.777.216	6.183.533	36,86%
/31	33.554.432	11.863.283	35,36%
/30	67.108.864	22.760.044	33,92%
/29	134.217.728	43.665.787	32,53%
/28	268.435.456	83.774.045	31,21%
/27	536.870.912	160.722.871	29,94%
/26	1.073.741.824	308.351.367	28,72%
/25	2.147.483.648	591.580.804	27,55%
/24	4.294.967.296	1.134.964.479	26,43%
/23	8.589.934.592	2.177.461.403	25,35%
/22	17.179.869.184	4.177.521.189	24,32%
/21	34.359.738.368	8.014.692.369	23,33%
/20	68.719.476.736	15.376.413.635	22,38%

78

Esta tabla presenta el porcentaje de utilización de bloques /56 en base al cálculo de HD-Ratio igual a 0,94.

Proveedores

- NTT Communications
 - Japón
 - IPv6 nativo (ADSL)
 - /48 a usuarios finales
 - http://www.ntt.com/business_e/service/category/nw_ipv6.html
- Internode
 - Australia
 - IPv6 nativo (ADSL)
 - /64 dinámico para sesiones PPP
 - Delega /60 fijos
 - <http://ipv6.internode.on.net/configuration/adsl-faq-guide/>

79

En cuanto a la política seguida por algunos proveedores a nivel mundial que ya proporcionan direcciones IPv6 a sus clientes, hay varios puntos de vista diferentes. Consideremos los siguientes ejemplos:

- NTT Communications
 - Japón
 - IPv6 nativo (ADSL)
 - /48 a usuarios finales
 - http://www.ntt.com/business_e/service/category/nw_ipv6.html
- Internode
 - Australia
 - IPv6 nativo (ADSL)
 - /64 dinámico para sesiones PPP
 - Delega /60 fijos
 - <http://ipv6.internode.on.net/configuration/adsl-faq-guide/>
- IIJ
 - Japón
 - Túneles
 - /48 a usuarios finales
 - <http://www.iiij.ad.jp/en/service/IPv6/index.html>
- Arcnet6
 - Malasia
 - IPv6 nativo (ADSL) o túneles
 - /48 a usuarios finales
 - se pueden distribuir bloques /40 y /44 (sujeto a aprobación)
 - <http://arcnet6.net.my/how.html>

Provedores

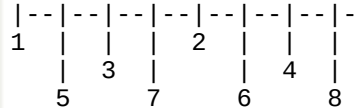
- Iij
 - Japón
 - Túneles
 - /48 a usuarios finales
 - <http://www.ij.ad.jp/en/service/IPv6/index.html>
- Arcnet6
 - Malasia
 - IPv6 nativo (ADSL) o túneles
 - /48 a usuarios finales
 - se pueden distribuir bloques /40 y /44 (sujeto a aprobación)
 - <http://arcnet6.net.my/how.html>

Consideraciones

- /32 =
 - 65 mil redes /48 (33 mil si consideramos las direcciones que se desperdician)
 - 16 millones de redes /56 (6 millones si consideramos el HD-ratio)
 - ¿Es suficiente para su proveedor?
 - Reservar un bloque (/48 ?) para infraestructura...

- Enlaces punto a punto:
 - /64? /112? /120? /126? /127?

- RFC 3531



81

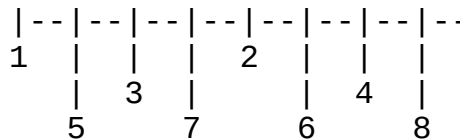
egi.br

Antes de solicitar un bloque de direcciones IPv6 se debe elaborar cuidadosamente el plan de direccionamiento de la red. Hay que considerar algunos aspectos importantes:

Un bloque /32 permite 65 mil redes /48, o 33 mil si consideramos las direcciones que se desperdician; 16 millones de redes /56, o 6 millones si consideramos el HD-ratio. Cada proveedor deberá analizar si estos valores satisfacen sus necesidades o si será necesario solicitar al Registro de Internet un bloque de direcciones más grande.

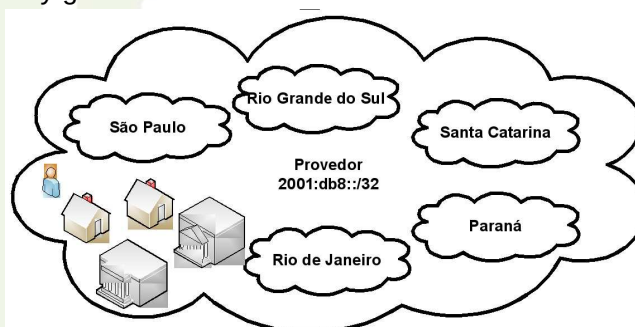
También se deben considerar las direcciones que se utilizarán para infraestructura, qué tamaño de bloque se reservará para las redes internas, para enlaces punto a punto, etc.

La RFC 3531 propone un método para administrar la asignación de los bits de un bloque de direcciones IPv6 o de un prefijo, de modo que éstos se distribuyan reservando un espacio mayor entre sí. De este modo, si fuera necesario distribuir nuevos bloques a un usuario, la probabilidad de distribuir bloques contiguos a los que dicho usuario posee será mayor.



Ejercicio de direccionamiento IPv6

- Usted es un proveedor y recibe un bloque **2001:0db8::/32**
- Usted tiene presencia en varias localidades (5 estados diferentes) y además tiene planes de expansión.
- Usted atiende a usuarios residenciales y a empresas pequeñas, medianas y grandes.



82

cgi.br

A continuación proponemos un ejercicio rápido para ver cómo planificar el direccionamiento:

- Usted es un proveedor y recibe un bloque **2001:0db8::/32**
- Usted tiene presencia en varias localidades (5 estados diferentes) y además tiene planes de expansión.
- Usted atiende a usuarios residenciales y a empresas pequeñas, medianas y grandes.

Trabaje los siguientes puntos:

- (1) Se ha decidido que la mejor manera de dividir las direcciones es jerárquicamente... ¿Cuál es el tamaño de bloque en cada estado?
- (2) ¿Qué tamaño de bloque se asignará a cada tipo de usuario?
- (3) ¿Cuántos usuarios de cada tipo se podrán atender de esta forma?
- (4) Indique el bloque correspondiente a cada localidad.
- (5) Escoja una localidad e indique los bloques correspondientes a cada tipo de usuario
- (6) En la misma localidad, indique el primer bloque y el segundo bloque asignados a cada tipo de usuario (los 2 primeros usuarios de cada tipo)
- (7) Indique la primera y la última dirección para el segundo bloque/usuario de cada tipo.

Ejercicio de direccionamiento IPv6

- (1) Se ha decidido que la mejor manera de dividir las direcciones es jerárquicamente... ¿Cuál es el tamaño de bloque en cada estado?
- (2) ¿Qué tamaño de bloque se asignará a cada tipo de usuario?
- (3) ¿Cuántos usuarios de cada tipo se podrán atender de este modo?
- (4) Indique el bloque correspondiente a cada localidad.
- (5) Escoja una localidad e indique los bloques correspondientes a cada tipo de usuario
- (6) En esa misma localidad, indique el primer bloque y el segundo bloque asignados a cada tipo de usuario (los 2 primeros usuarios de cada tipo)
- (7) Indique la primera y la última dirección para el segundo bloque/usuario de cada tipo.

Más información:

- RFC 5375 - *IPv6 Unicast Address Assignment Considerations*
- RFC 3177 - IAB/IESG Recommendations on IPv6 Address Allocations to Sites
- RFC 3531 - *A Flexible Method for Managing the Assignment of Bits of an IPv6 Address Block*
- RFC 3627 - *Use of /127 Prefix Length Between Routers Considered Harmful*
- RFC 3194 - *The Host-Density Ratio for Address Assignment Efficiency: An update on the HD ratio*
- RFC 4692 - *Considerations on the IPv6 Host Density Metric*
- <http://www.potaroo.net/ispcol/2005-07/ipv6size.html>
- <http://www.lacnic.net/en/politicas/manual12.html>
- <http://tools.ietf.org/html/draft-narten-ipv6-3177bis-48boundary-04>
- https://www.arin.net/policy/proposals/2005_8.html
- <http://www.apnic.net/policy/ipv6-address-policy#2.7>
- <http://www.ripe.net/ripe/docs/ipv6-sparse.html>
- <http://www.ipv6book.ca/allocation.html>
- <http://tools.ietf.org/html/draft-kohno-ipv6-prefixlen-p2p-00>
- http://www.swinog.ch/meetings/swinog18/swissix_swinog_18.pdf
- https://www.arin.net/participate/meetings/reports/ARIN_XXIV/PDF/wednesday/ipv6_implementation_fundamentals.pdf
- http://www.6deploy.org/workshops/20090921_bogota_colombia/Consulintel_IPv6_3-Direccionamiento_IPv6.pdf

IPv6.br

La nueva generación del
Protocolo de Internet

Funcionalidades de IPv6 #1

Módulo 4

86

El protocolo IPv6 presenta una serie de nuevas funcionalidades y otras mejoras con respecto a IPv4.

En la primera parte de este módulo conoceremos un poco más sobre estas funcionalidades, comenzando por el estudio del protocolo ICMPv6 (*Internet Control Message Protocol versión 6*), pieza fundamental para la ejecución de herramientas tales como el protocolo de Descubrimiento de Vecinos (*Neighbor Discovery*) y los mecanismos de autoconfiguración *stateless*, que también se tratan aquí. Analizaremos el funcionamiento del protocolo DHCPv6 y cómo estas funcionalidades pueden contribuir al trabajo de renumeración de redes.

ICMPv6

- Definido en la RFC 4443.
- Mismas funciones que ICMPv4 (pero no compatibles):
 - Informar características de la red;
 - Realizar diagnósticos;
 - Informar errores en el procesamiento de paquetes.
- Asume las funcionalidades de otros protocolos:
 - ARP/RARP
 - IGMP
- Identificado por el valor 58 en el campo Siguiente Encabezado.
- Se debe implementar en todos los nodos.

Definido en la RFC 4443 para ser utilizado con IPv6, ICMPv6 es una versión actualizada del ICMP (*Internet Control Message Protocol*) que se utiliza con IPv4.

Aunque tiene las mismas funciones que ICMPv4 (como por ejemplo informar errores en el procesamiento de paquetes y reenviar mensajes sobre el estado o las características de la red), esta nueva versión del ICMP no es compatible con su predecesor y ahora presenta un mayor número de mensajes y funcionalidades.

Ahora ICMPv6 es responsable de realizar las funciones de los protocolos ARP (*Address Resolution Protocol*), que mapean las direcciones de capa dos a las IP y viceversa en IPv4, y del IGMP (*Internet Group Management Protocol*), que gestiona los miembros de los grupos *multicast* en IPv4.

El valor del campo Siguiente Encabezado (que indica la presencia del protocolo ICMPv6) es 58, y en todos los nodos se debe implementar soporte para este protocolo.

ICMPv6

- Es precedido por los encabezados de extensión (si los hubiera) y por el encabezado base de IPv6.

IPv6
cadena de encab. de extensión
ICMPv6

- Protocolo clave de la arquitectura IPv6.
- Esencial para las funcionalidades de IPv6:
 - Gestión de grupos *multicast*;
 - Descubrimiento de Vecinos (*Neighbor Discovery*);
 - Movilidad IPv6;
 - Descubrimiento de la *Path MTU*.

En un paquete IPv6, el ICMPv6 se coloca inmediatamente después del encabezado base de IPv6 y de los encabezados de extensión (si los hubiera).

ICMPv6, es un protocolo clave en la arquitectura IPv6 ya que, además de gestionar los grupos *multicast*, a través del protocolo MLD (*Multicast Listener Discovery*) y de la resolución de direcciones de capa dos, sus mensajes son esenciales para el funcionamiento del protocolo de Descubrimiento de Vecinos (*Neighbor Discovery*), responsable de localizar routers vecinos en la red, detectar cambios de dirección en el enlace, detectar direcciones duplicadas, etc.; para el soporte a la movilidad, gestionando las Direcciones de Origen de los *hosts* dinámicamente; y para el proceso de descubrimiento de la menor MTU (*Maximum Transmit Unit*) en el camino de un paquete hasta su destino.

ICMPv6

- Encabezados simples

Tipo (Type)	Código (Code)	Soma de Verificação (Checksum)
Dados		

- **Tipo** (8 bits): Especifica el tipo de mensaje.
- **Código** (8 bits): Ofrece algunos datos adicionales para determinados tipos de mensajes.
- **Suma de Verificación** (16 bits): Se utiliza para detectar datos corruptos en el encabezado ICMPv6 y en parte del encabezado IPv6.
- **Datos**: Presenta información de diagnóstico y error según el tipo de mensaje. Su tamaño puede variar dependiendo del mensaje.

89

El encabezado de todos los mensajes ICMPv6 tienen la misma estructura y está compuesto por cuatro campos:

- **Tipo**: Especifica el tipo de mensaje, lo que determinará el formato del cuerpo del mensaje Su tamaño es de ocho bits;
- **Código**: Ofrece algunos datos adicionales para determinados tipos de mensajes. Su tamaño también es de ocho bits;
- **Suma de Verificación**: Se utiliza para detectar datos corruptos en el encabezado ICMPv6 y en parte del encabezado IPv6. Su tamaño es de 16 bits;
- **Datos**: Presenta información de diagnóstico y error según el tipo de mensaje. Para ayudar en la resolución de problemas, los mensajes de error incluyen en este campo el paquete que invocó el mensaje, ya que el tamaño total del paquete ICMPv6 no debe exceder la mínima MTU de IPv6, que es de 1280 bytes.

ICMPv6

- Tiene dos clases de mensajes:
 - Mensajes de error
 - *Destination Unreachable*
 - *Packet Too Big*
 - *Time Exceeded*
 - *Parameter Problem*
 - Mensajes de información
 - *Echo Request* y *Echo Reply*
 - *Multicast Listener Query*
 - *Multicast Listener Report*
 - *Multicast Listener Done*
 - *Router Solicitation* y *Router Advertisement*
 - *Neighbor Solicitation* y *Neighbor Advertisement*
 - *Redirect...*

90

Los mensajes ICMPv6 se dividen en dos clases, cada una de ellas compuesta por diferentes tipos de mensajes, de acuerdo con las siguientes tablas:

Mensajes de error:

Tipo	Nombre	Descripción
1	Destination Unreachable	Indica fallas en la entrega del paquete (como dirección o puerto desconocido) o problemas en la comunicación.
2	Packet Too Big	Indica que el tamaño del paquete es mayor que la Unidad Máxima de Transferencia (MTU) de un enlace.
3	Time Exceeded	Indica que el Límite de Direccionamiento o el tiempo de ensamble del paquete fue excedido.
4	Parameter Problem	Indica un error en alguno de los campos del encabezado IPv6 o que el tipo indicado en el campo Siguiente Encabezado no fue reconocido.
100-101		Uso experimental
102-126		No utilizados
127		Reservado para la expansión de mensajes de error ICMPv6

Mensajes de información

Tipo	Nombre	Descripción
128	Echo Request	Utilizados por el comando ping.
129	Echo Reply	
130	Multicast Listener Query	Utilizados en la gestión de grupos <i>multicast</i> .
131	Multicast Listener Report	
132	Multicast Listener Done	
133	Router Solicitation	Utilizados con el protocolo de Descubrimiento de Vecinos.
134	Router Advertisement	
135	Neighbor Solicitation	
136	Neighbor Advertisement	
137	Redirect Message	
138	Router Renumbering	Utilizado en el mecanismo de re-direccionamiento(<i>Renumbering</i>) de routers.
139	ICMP Node Information Query	Utilizados para descubrir datos sobre nombres y direcciones, actualmente están limitados a herramientas de diagnóstico, depuración y gestión de redes.
140	ICMP Node Information Response	
141	Inverse ND Solicitation Message	Utilizados en una extensión del protocolo de Descubrimiento de Vecinos.
142	Inverse ND Advertisement Message	
143	Version 2 Multicast Listener Report	Utilizado en la gestión de grupos <i>multicast</i> .
144	HA Address Discovery Req. Message	Utilizados en el mecanismo de Movilidad IPv6.
145	HA Address Discovery Reply Message	
146	Mobile Prefix Solicitation	
147	Mobile Prefix Advertisement	
148	Certification Path Solicitation Message	Utilizados por el protocolo SEND.
149	Cert. Path Advertisement Message	
150		Utilizado experimentalmente con protocolos de movilidad como <i>Seamoby</i> .
151	Multicast Router Advertisement	Utilizados por el mecanismo <i>Multicast Router Discovery</i>
152	Multicast Router Solicitation	
153	Multicast Router Termination	
154	FMIPv6 Messages	Utilizado por el protocolo de movilidad <i>Fast Handovers</i>
200-201		Uso experimental
255		Reservado para la expansión de mensajes de error ICMPv6

Descubrimiento de Vecinos

- *Neighbor Discovery* – definido en la RFC 4861.
- Asume las funciones de los ARP, *ICMP Router Discovery* y *ICMP Redirect* de IPv4.
- Agrega nuevos métodos que no existían en la versión anterior del protocolo IP.
- Agiliza algunos procesos de configuración de red:
 - determinar la dirección MAC de los nodos de la red;
 - encontrar routers vecinos;
 - determinar prefijos y otros datos de configuración de la red;
 - detectar direcciones duplicadas;
 - determinar la accesibilidad de los routers;
 - redireccionamiento de paquetes;
 - autoconfiguración de direcciones.

92

Definido por la RFC4861, el protocolo de Descubrimiento de Vecinos agiliza algunos procesos de configuración de red con respecto a IPv4 combinando las funciones de protocolos como ARP, *ICMP Router Discovery* y *ICMP Redirect* y además agrega nuevos métodos que no existían en la versión anterior del protocolo IP.

El protocolo de Descubrimiento de Vecinos de IPv6 es utilizado por los *hosts* y routers para los siguientes propósitos:

- determinar la dirección MAC de los nodos de la red;
- encontrar routers vecinos;
- determinar prefijos y otros datos de configuración de la red;
- detectar direcciones duplicadas;
- determinar la accesibilidad de los routers;
- redireccionamiento de paquetes;
- autoconfiguración de direcciones.

Descubrimiento de Vecinos

- Utiliza 5 tipos de mensajes ICMPv6:
 - *Router Solicitation* (RS) – ICMPv6 Tipo 133;
 - *Router Advertisement* (RA) – ICMPv6 Tipo 134;
 - *Neighbor Solicitation* (NS) – ICMPv6 Tipo 135;
 - *Neighbor Advertisement* (NA) – ICMPv6 Tipo 136;
 - *Redirect* – ICMPv6 Tipo 137.
- Se configuran con el valor 255 en el campo Límite de Direccionamiento.
- Pueden o no contener opciones:
 - *Source link-layer address*.
 - *Target link-layer address*.
 - *Prefix information*.
 - *Redirected header*.
 - MTU.

93

Los mensajes *Neighbor Discovery* se configuran con un Límite de Direccionamiento de 255 para asegurar que los mensajes recibidos tengan su origen en un nodo del mismo enlace, descartando los mensajes que tengan valores diferentes.

Neighbor Discovery utiliza cinco mensajes ICMPv6:

- **Router Solicitation (ICMPv6 tipo 133):** Utilizado por los *hosts* para solicitar a los routers mensajes *Router Advertisements* inmediatamente. Normalmente se envía a la dirección **multicast FF02::2** (*all-routers on link*);
- **Router Advertisement (ICMPv6 tipo 134):** Se envía periódicamente o en respuesta a un *Router Solicitation* y es utilizado por los routers para anunciar su presencia en un enlace. Los mensajes periódicos se envían a la dirección **multicast FF02::1** (*all-nodes on link*), mientras que los solicitados se envían directamente a la dirección del solicitante. Un RA transporta información referente a la configuración de la red, por ejemplo:
 - El valor por defecto del enlace para el campo Límite de Direccionamiento;
 - Un *flag* que especifica si se debe utilizar autoconfiguración *stateless* o *stateful*;
 - Otro *flag* que especifica si los nodos deben utilizar configuraciones *stateful* para obtener otros datos acerca de la red;
 - En las redes que soportan movilidad IPv6 se utiliza un tercer *flag* para indicar si el router es un Agente de Origen;

- Por cuánto tiempo (en segundos) el router será considerado el router por defecto del enlace. Si no es el router por defecto este valor será cero;
- El tiempo que un *host* supone que los vecinos son alcanzables luego de recibir una confirmación de accesibilidad;
- El intervalo entre el envío de mensajes *Neighbor Solicitation*.
- ***Neighbor Solicitation (ICMPv6 tipo 135)***: Mensaje *multicast* enviado por un nodo para determinar la dirección MAC y la accesibilidad de un vecino y detectar la existencia de direcciones duplicadas. Este mensaje tiene un campo para indicar la dirección de origen del mensaje;
- ***Neighbor Advertisement (ICMPv6 tipo 136)***: Se envía como respuesta a un *Neighbor Solicitation*, pero también se puede enviar para anunciar el cambio de alguna dirección dentro del enlace. Este mensaje tiene tres *flags*:
 - La primera indica si quien está enviando el mensaje es un router;
 - La segunda indica si el mensaje es una respuesta a un NS;
 - La tercera indica si la información que transporta el mensaje es una actualización de la dirección de alguno de los nodos de la red.
- ***Redirect (ICMPv6 tipo 137)***: Utilizado por los routers para informar al *host* el mejor router para encaminar el paquete a su destino. Este mensaje contiene información como la dirección del router que se considera el mejor salto y la dirección del nodo que está siendo redireccionado.

Estos mensajes pueden tener o no opciones, definidas en la RFC 4861:

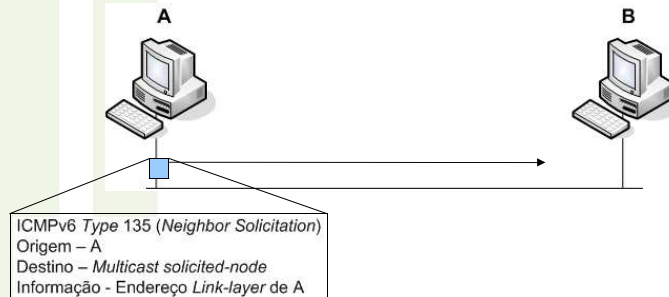
- ***Source link-layer address***: Contiene la dirección MAC del remitente del paquete. Se utiliza en los mensajes NS, RS y RA;
- ***Target link-layer address***: Contiene la dirección MAC del destino del paquete. Se utiliza en los mensajes NA y *Redirect*;
- ***Prefix information***: Proporciona a los *hosts* los prefijos del enlace y los prefijos para autoconfiguración de direcciones. Se utiliza en los mensajes RA;
- ***Redirected header***: Contiene todo el paquete que está siendo redireccionado o parte del mismo. Se utiliza en los mensajes *Redirect*; y
- ***MTU***: Indica el valor de la MTU del enlace. Se utiliza en los mensajes RA.

Se definieron nuevas opciones para las nuevas funcionalidades del protocolo de Descubrimiento de Vecinos. Estas opciones se describirán a medida que presentemos estas nuevas funciones.

Descubrimiento de Vecinos

- **Descubrimiento de direcciones de capa de enlace**

- Determina la dirección MAC de los vecinos del mismo enlace.
- Reemplaza al protocolo ARP.
- Utiliza la dirección *multicast solicited-node* en lugar de la dirección *broadcast*.
- El *host* envía un mensaje NS informando su dirección MAC y solicita la dirección MAC del vecino.



95

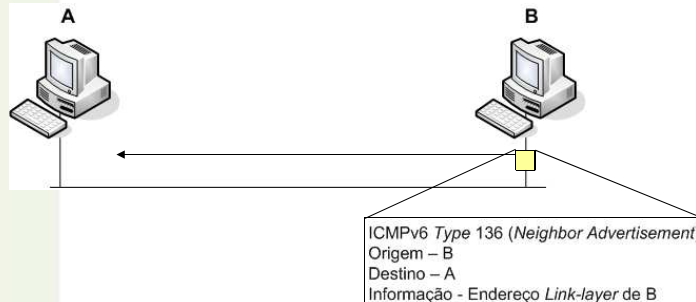
egi.br

Esta funcionalidad se utiliza para determinar la dirección MAC de los vecinos del mismo enlace: un *host* envía un mensaje NS a la dirección *multicast solicited node* del vecino informando su dirección MAC.

Descubrimiento de Vecinos

- **Descubrimiento de direcciones de capa de enlace**

- Determina la dirección MAC de los vecinos del mismo enlace.
- Reemplaza al protocolo ARP.
- Utiliza la dirección *multicast solicited-node* en lugar de la dirección *broadcast*.
- El *host* envía un mensaje NS informando su dirección MAC y solicita la dirección MAC del vecino.
- El vecino responde enviando un mensaje NA informando su dirección MAC.



96

egi.br

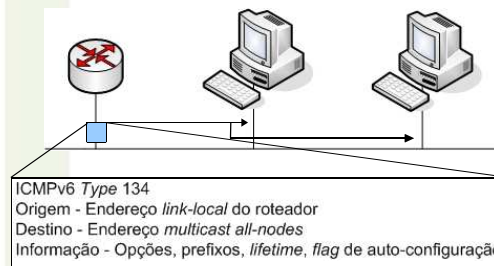
Al recibir el mensaje, el vecino responde enviando un mensaje NA informando su dirección MAC.

En IPv6, esta característica del protocolo de Descubrimiento de Vecinos reemplaza al protocolo ARP de IPv4 y, en lugar de utilizar una dirección *broadcast*, utiliza la dirección *multicast solicited-node* como dirección de destino.

Descubrimiento de Vecinos

- **Descubrimiento de routers y prefijos**

- Localizar routers vecinos dentro del mismo enlace.
- Determina prefijos y parámetros relacionados con la autoconfiguración de direcciones.
- En IPv4 esta función es realizada por los mensajes *ARP Request*.
- Los routers envían mensajes RA a la dirección *multicast all-nodes*.



97

Esta funcionalidad del protocolo de Descubrimiento de Vecinos se utiliza para localizar routers vecinos dentro del mismo enlace, así como para aprender prefijos y parámetros relacionados con la autoconfiguración de direcciones.

Esta información se envía desde un router local, a través de mensajes RA encaminados a la dirección *multicast all-nodes*.

En IPv4 el mapeo de las direcciones de la red local se realiza a través de mensajes *ARP Request*.

Descubrimiento de Vecinos

- **Detección de direcciones duplicadas**

- Verifica la unicidad de las direcciones de un nodo dentro del enlace.
- Se debe realizar antes de atribuir cualquier dirección *unicast* a una interfaz.
- Consiste en el envío de un mensaje NS por parte del *host* con su propia dirección en el campo *target address*. Si como respuesta se recibe un mensaje NA, esto indica que la dirección ya está siendo utilizada.

La Detección de Direcciones Duplicadas es el procedimiento que utilizan los nodos para verificar la unicidad de las direcciones en un enlace y se debe realizar antes de atribuir cualquier dirección *unicast* a una interfaz, sin importar si éstas se obtienen mediante autoconfiguración *stateless*, DHCPv6 o configuración manual.

Este mecanismo consiste en el envío de un mensaje NS por parte del *host* con su propia dirección en el campo *target address*. Si como respuesta se recibe un mensaje NA, esto indica que la dirección ya está siendo utilizada y que el proceso de configuración debe ser interrumpido.

En IPv4, los nodos utilizan mensajes *ARP Request* y el método llamado *gratuitous ARP* para detectar direcciones *unicast* duplicadas dentro del mismo enlace, definiendo los campos *Source Protocol Address* y *Target Protocol Address*, del encabezado del mensaje *ARP Request*, con la dirección IPv4 que está siendo verificada.

Descubrimiento de Vecinos

- **Detección de vecinos inaccesibles**

- Se utiliza para rastrear la accesibilidad de los nodos a lo largo del camino.
- Un nodo considera que un vecino es accesible si recientemente ha recibido confirmación de la entrega de algún paquete a dicho vecino.
 - Puede ser una respuesta a mensajes del protocolo de Descubrimiento de Vecinos o algún proceso de capa de transporte que indique que se estableció una conexión.
- Se aplica solamente para direcciones *unicast*.
- *Neighbor Cache* (similar a la tabla ARP).
- *Destination Cache*.

99

Este mecanismo se utiliza en comunicaciones *host-a-host*, *host-a-router* y *router-a-host* para rastrear la accesibilidad de los nodos a lo largo del camino.

Un nodo considera que un vecino es accesible si recientemente ha recibido confirmación de la entrega de algún paquete a dicho vecino. Esta confirmación se puede dar de dos maneras diferentes: una respuesta a un mensaje del protocolo de Descubrimiento de Vecinos o algún proceso de capa de transporte que indique que se estableció una conexión.

Este proceso solo se ejecuta cuando se envían paquetes a una dirección *unicast*; no se utiliza cuando se envían paquetes a direcciones *multicast*.

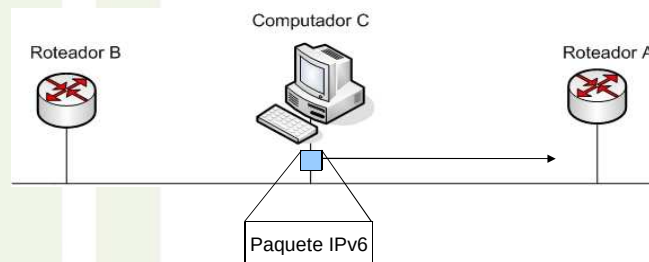
Para realizar el seguimiento de los estados de un vecino el nodo IPv6 utiliza dos tablas principales:

- ***Neighbor Cache*** – Mantiene una lista de vecinos locales a los cuales se envió tráfico recientemente, almacenando sus direcciones IP, información sobre la dirección MAC y un *flag* que indica si el vecino es un router o un *host*. También informa si todavía hay paquetes en cola esperando ser enviados, la accesibilidad de los vecinos y la próxima vez que está programado un evento de detección de vecinos inaccesibles. Esta tabla es comparable a la tabla ARP de IPv4.
- ***Destination Cache*** – Mantiene información sobre destinos a los cuales se envió tráfico recientemente (tanto destinos locales como remotos) y se actualiza con información recibida a través de mensajes *Redirect*. El *Neighbor Cache* se puede considerar como un subconjunto de la información del *Destination Cache*.

Descubrimiento de Vecinos

• Redireccionamiento

- Envía mensajes *Redirect*
- Redirecciona un *host* a un router más apropiado para el primer salto.
- Informar al *host* qué destino se encuentra en el mismo enlace.
- Este mecanismo es igual al que existe en IPv4.



100

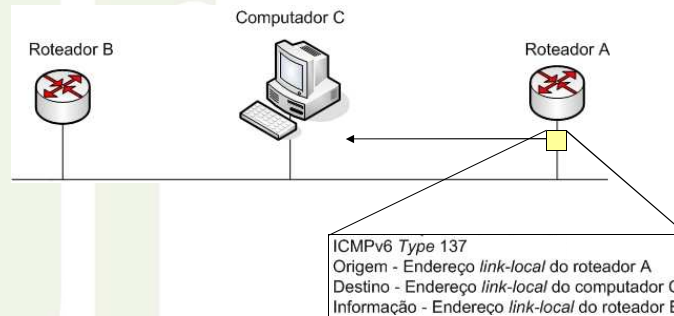
cgi.br

Los routers envían mensajes *Redirect* para redireccionar un *host* automáticamente a un router más apropiado o informar al *host* qué destino se encuentra en el mismo enlace. Este mecanismo es igual al que existe en IPv4.

Descubrimiento de Vecinos

- **Redireccionamiento**

- Envía mensajes *Redirect*
- Redirecciona un *host* a un router más apropiado para el primer salto.
- Informar al *host* qué destino se encuentra en el mismo enlace.
- Este mecanismo es igual al que existe en IPv4.

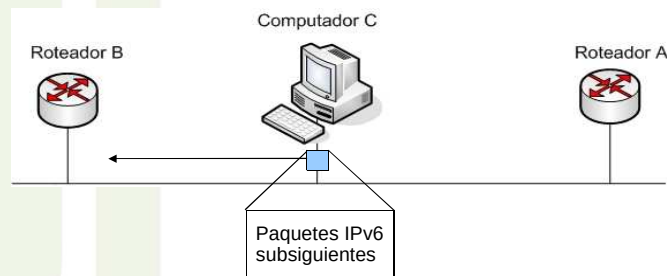


101

Descubrimiento de Vecinos

- **Redireccionamiento**

- Envía mensajes *Redirect*
- Redirecciona un *host* a un router más apropiado para el primer salto.
- Informar al *host* qué destino se encuentra en el mismo enlace.
- Este mecanismo es igual al que existe en IPv4.



102

Descubrimiento de Vecinos

- **Autoconfiguración de direcciones stateless**

- Mecanismo que permite atribuir direcciones *unicast* a los nodos...
 - sin necesidad de realizar configuraciones manuales.
 - sin utilizar servidores adicionales.
 - con una configuración mínima de los routers.
- Genera direcciones IP a partir de información enviada por los routers y datos locales como la dirección MAC.
- Genera una dirección para cada prefijo informado en los mensajes RA.
- Si no hay routers presentes en la red solamente se genera una dirección *link local*.
- Los routers solo usan este mecanismo para generar direcciones *link-local*.

103

El mecanismo de autoconfiguración *stateless*, definido en la RFC 4862, permite atribuir direcciones IPv6 a las interfaces sin necesidad de realizar configuraciones manuales, sin utilizar servidores adicionales (DHCP), apenas realizando una configuración mínima de los routers.

Para generar una dirección IP un *host* utiliza una combinación de datos locales, como la dirección MAC de la interfaz o un valor aleatorio para generar el ID, e información recibida de los routers, como múltiples prefijos. Si no hay routers presentes, el *host* solo genera la dirección *link local* con el prefijo **FE80::**.

Los routers solo usan este mecanismo para generar direcciones *link-local*. Sus direcciones globales se deben configurar de otra manera.

Descubrimiento de Vecinos

• *Autoconfiguración de Direcciones Stateless*

- Se genera una dirección *link-local*.
 - Prefijo **FE80::/64** + identificador de la interfaz.
- Dirección que se agrega a los grupos *multicast solicited-node* y *all-node*.
- Se verifica la unicidad de la dirección.
 - Si la dirección ya está siendo utilizada el proceso se interrumpe, requiriéndose una configuración manual.
 - Si la dirección es considerada única y válida ésta se atribuye a la interfaz.
- El *host* envía un mensaje RS al grupo *multicast all-routers*.
- Todos los routers del enlace responden con mensajes RA.
- Estados de las direcciones:
 - Dirección tentativa;
 - Dirección preferida;
 - Dirección desaprobada;
 - Dirección válida;
 - Dirección inválida.

104

El mecanismo de autoconfiguración de direcciones se ejecuta respetando los siguientes pasos:

- Se genera una dirección *link-local* agregando el prefijo **FE80::/64** al identificador de la interfaz;
- Esta dirección pasa a formar parte de los grupos *multicast solicited-node* y *all-node*;
- Se realiza la verificación de la unicidad de la dirección *link-local* generada;
 - ♦ Si otro nodo del enlace ya está utilizando la misma dirección el proceso de autoconfiguración se interrumpe y es necesario realizar una configuración manual;
- Si la dirección se considera única y válida ésta será automáticamente inicializada para la interfaz;
- El *host* envía un mensaje *Router Solicitation* al grupo *multicast all-routers*;
- Todos los routers del enlace responden con un mensaje *Router Advertisement* informando: los routers por defecto; un valor predefinido para el campo Límite de Direccionamiento; la MTU del enlace; la lista de prefijos de la red, para los cuales también se generarán direcciones automáticamente.

Una dirección IPv6 puede asumir diferentes estados:

- Dirección tentativa – Dirección que aun no ha sido atribuida. Es el estado previo a la atribución, mientras se realiza el proceso DAD. No se puede utilizar en las comunicaciones del nodo, sino solamente para los mensajes relacionados con el Descubrimiento de Vecinos;
- Dirección preferida – Dirección atribuida a la interfaz que puede ser utilizada sin restricciones hasta que expire su tiempo de vida;
- Dirección desaprobadada – Dirección cuyo tiempo de vida ha expirado. Se puede utilizar para continuar las comunicaciones abiertas por la misma, pero no para iniciar nuevas comunicaciones;
- Dirección válida – Término utilizado para designar tanto a las direcciones preferidas como a las direcciones desaprobadadas;
- Dirección inválida – Dirección que no se puede atribuir a una interfaz. Una dirección se vuelve inválida cuando expira su tiempo de vida.

DHCPv6

- **Autoconfiguración de direcciones stateful**

- Utilizada por el sistema cuando no se encuentra ningún router.
- Utilizada por el sistema cuando así lo indican los mensajes RA.
- Provee:
 - Direcciones IPv6
 - Otros parámetros (servidores DNS, NTP...)
- Los clientes utilizan una dirección *link-local* para transmitir o recibir mensajes DHCP.
- Los servidores utilizan direcciones *multicast* para recibir mensajes de los clientes (**FF02::1:2** o **FF05::1:3**).
- Los clientes envían mensajes a servidores fuera de su enlace utilizando un *Relay* DHCP.

106

El *Dynamic Host Configuration Protocol* (DHCP) es un protocolo de autoconfiguración *stateful* que se utiliza para distribuir dinámicamente direcciones IP en una red a partir de un servidor DHCP, permitiendo un mayor control de la atribución de direcciones a los *host*.

Definido en la RFC 3315, el protocolo DHCPv6 es una alternativa al mecanismo de autoconfiguración *stateless* de IPv6 que se puede utilizar cuando no hay routers en la red o cuando su uso es indicado en los mensajes RA; este mecanismo puede proveer direcciones IPv6 y diferentes parámetros de la red, como por ejemplo direcciones de servidores DNS, NTP, SIP, etc.

En DHCPv6 el intercambio de mensajes entre cliente y servidor se realiza usando el protocolo UDP. Los clientes utilizan una dirección *link-local* para transmitir o recibir mensajes DHCP, mientras que los servidores utilizan una dirección *multicast* reservada (**FF02::1:2** o **FF05::1:3**) para recibir mensajes de los clientes. Cuando el cliente necesita enviar un mensaje a un servidor que está fuera de su subred se utiliza un *Relay* DHCP.

DHCPv6

- **Autoconfiguración de direcciones stateful**

- Permite un mayor control de la atribución de direcciones a los *host*.
- Los mecanismos de autoconfiguración de direcciones *stateful* y *stateless* se pueden utilizar simultáneamente.
 - Por ejemplo: Utilizar autoconfiguración *stateless* para atribuir las direcciones y DHCPv6 para informar la dirección del servidor DNS.
- DHCPv6 y DHCPv4 son independientes. Las redes con doble pila requieren servicios DHCP separados.

107

cgi.br

El uso de DHCPv6 permite un mayor control de la atribución de direcciones ya que, además de proporcionar opciones de configuración de red, permite definir políticas para la distribución de direcciones y atribuir direcciones a los *hosts* que no se derivan de la dirección MAC.

En una red IPv6 se puede combinar el uso de autoconfiguración *stateless* con servidores DHCP. En este escenario es posible, por ejemplo, utilizar autoconfiguración *stateless* para atribuir direcciones a los *hosts* y servidores DHCPv6 para proveer información de configuración adicional, como por ejemplo la dirección de los servidores DNS.

Los protocolos DHCPv6 y DHCPv4 son independientes, de modo que en una red con doble pila se debe ejecutar un servicio para cada protocolo. En el caso de DHCPv4 es necesario configurar en el cliente si éste utilizará DHCP, mientras que la utilización de DHCPv6 se indica a través de las opciones de los mensajes RA.

Renumeración de la red

- *Hosts* – Autoconfiguración *stateless* o DHCPv6
- *Routers* – *Router Renumbering*
- Mensajes ICMPv6 Tipo 138
- Formato de los mensajes
 - Encabezado RR + Cuerpo del mensaje

Tipo	Código	Suma de verificación
Número secuencial		
Número de segmento	Flags	Retraso máximo
Reservado		
Cuerpo del mensaje Mensaje de comando / Mensaje de resultado		

108

Muchas veces el direccionamiento de una red se basa en los prefijos atribuidos por los ISP. Si se cambia de proveedor es necesario reenumerar todas las direcciones de la red.

En IPv6 el proceso de redireccionamiento de los *host* se puede realizar de manera relativamente sencilla. A través de los mecanismos del protocolo de Descubrimiento de Vecinos, el router puede anunciar un nuevo prefijo a todos los *hosts* del enlace. También es posible utilizar servidores DHCPv6. Para tratar la configuración y reconfiguración de prefijos en los routers tan fácilmente como en los *hosts*, en la RFC 2894 se definió el protocolo *Router Renumbering*.

El mecanismo *Router Renumbering* utiliza mensajes ICMPv6 de tipo 138, los cuales se envían a los routers a través de la dirección *multicast all-routers*, y contienen las instrucciones para actualizar sus prefijos.

Los mensajes *Router Renumbering* están formados por los siguientes campos:

- Tipo - 138 (decimal);
- Código - 0 para mensajes de comando;
 - 1 para mensajes de resultado;
 - 255 para poner en cero el número secuencial;

- Suma de Verificación – Verifica la integridad de los mensajes ICMPv6 y de parte del encabezado IPv6;
- Número secuencial – Identifica las operaciones;
- Número de segmento – Enumera diferentes mensajes RR válidos que tienen el mismo número secuencial.
- *Flags* – T: Indica si la configuración del router debe ser modificada o si se trata de una prueba;
R: Indica si se debe enviar un mensaje de resultado;
A: Indica si el comando se debe aplicar a todas las interfaces, independientemente de su estado;
S: Indica que el comando se debe aplicar a todas las interfaces, independientemente de la subred a la cual pertenezcan;
P: Indica que el mensaje de resultado contiene el informe completo del procesamiento del mensaje de comando, o que el mensaje de comando fue tratado previamente (y no se trata de una prueba) y que el router no lo está procesando nuevamente.
- Retraso máximo – Especifica el tiempo máximo, en milisegundos, que un router debe retrasar el envío de cualquier respuesta a un mensaje de comando.

Los mensajes de comando están formados por secuencias de operaciones, *Match-Prefix* y *Use-Prefix*. *Match-Prefix* indica cuál prefijo debe ser modificado, mientras que *Use-Prefix* indica el nuevo prefijo. Las operaciones pueden ser ADD, CHANGE o SET-GLOBAL, que indican, respectivamente, que el router debe agregar los prefijos indicados en *Use-Prefix* al conjunto de prefijos configurados; eliminar el prefijo indicado en *Match-Prefix* (si es que existen) y cambiarlos por los prefijos contenidos en *Use-Prefix*; o reemplazar todos los prefijos de alcance global por los prefijos de *Use-Prefix*. Si el conjunto *Use-Prefix* es un conjunto vacío, la operación ADD no realiza ninguna adición y las otras dos operaciones simplemente borran el contenido indicado.

Los routers también envían mensajes de resultados que contienen un *Match Report* para cada prefijo igual a los enviados en el mensaje de comando.

Más información:

- RFC 3315 - *Dynamic Host Configuration Protocol for IPv6* (DHCPv6)
- RFC 4443 - *Internet Control Message Protocol (ICMPv6) for the Internet Protocol Version 6 (IPv6) Specification*
- RFC 4861 - *Neighbor Discovery for IP version 6* (IPv6)
- RFC 5006 - *IPv6 Router Advertisement Option for DNS Configuration*

IPv6.br

La nueva generación del
Protocolo de Internet

110

Funcionalidades de IPv6 #2

Módulo 4

111

Continuando con el estudio de las funcionalidades del protocolo IPv6, a continuación veremos cómo el nuevo protocolo trata el tema de la fragmentación de paquetes y la gestión de grupos *multicast*. También veremos las modificaciones introducidas al protocolo DNS y las mejoras que presenta el protocolo IPv6 relacionadas con la aplicación de QoS y el soporte para movilidad.

Path MTU Discovery

- MTU - *Maximum Transmit Unit* – Tamaño máximo de paquete que puede atravesar el enlace.
- Fragmentación – Permite enviar paquetes mayores que la MTU de un enlace.
 - IPv4 – Todos los routers pueden fragmentar los paquetes mayores que la MTU del siguiente enlace.
 - Dependiendo del diseño de la red, un paquete IPv4 puede ser fragmentado más de una vez durante su trayecto.
 - IPv6 – La fragmentación se realiza solamente en el origen.
- *Path MTU Discovery* – Busca garantizar que el paquete encaminado sea del mayor tamaño posible.
- Todos los nodos IPv6 deben soportar PMTUD.
 - Las implementaciones mínimas de IPv6 pueden omitir este soporte, utilizando 1280 bytes como máximo tamaño de paquete.

Dependiendo de los protocolos de enrutamiento cada enlace de la red puede tener un valor de MTU diferente, es decir una limitación diferente respecto del tamaño máximo de paquete que puede atravesarlo. Para poder encaminar paquetes mayores que la MTU del enlace, éstos se deben fragmentar en paquetes menores que serán ensamblados al llegar a su destino.

En la transmisión paquetes IPv4 cada router a lo largo del camino puede fragmentar los paquetes si éstos son mayores que la MTU del siguiente enlace. Dependiendo del diseño de la red, un paquete IPv4 puede ser fragmentado más de una vez durante su trayecto y ensamblado nuevamente en el destino final.

En IPv6 la fragmentación de los paquetes se realiza solamente en el origen, no estando permitida en los routers intermedios. Este proceso tiene por objeto reducir el *overhead* del cálculo de los encabezados modificados en los routers intermedios.

Para ello en el inicio del proceso de fragmentación se utiliza el protocolo *Path MTU Discovery*, descrito en la RFC 1981, que descubre de forma dinámica cuál es el tamaño máximo de paquete permitido, identificando previamente las MTU de cada enlace en el camino hasta el destino. Todos los nodos IPv6 deben soportar el protocolo PMTUD. No obstante, las implementaciones mínimas de IPv6 pueden omitir este soporte, utilizando 1280 bytes como máximo tamaño de paquete.

Path MTU Discovery

- Asume que la máxima MTU del camino es igual a la MTU del primer salto.
- Los paquetes mayores que el soportado por algún router a lo largo del camino se descartan.
 - Se devuelve un mensaje ICMPv6 *packet too big*.
- Luego de recibir este mensaje el nodo de origen reduce el tamaño de los paquetes de acuerdo con la MTU indicada en el mensaje *packet too big*.
- El procedimiento finaliza cuando el tamaño del paquete es igual o menor que la menor MTU del camino.
- Estas iteraciones pueden ocurrir varias veces hasta encontrar la menor MTU.
- Los paquetes enviados a un grupo *multicast* utilizan un tamaño igual a la menor PMTU de todo el conjunto de destinos.

113

El proceso *Path MTU Discovery* comienza suponiendo que la MTU de todo el camino es igual a la MTU del primer salto. Si el tamaño de los paquetes enviados es mayor que el que soporta alguno de los routers a lo largo del camino, éste lo descargará y devolverá un mensaje ICMPv6 *packet too big*, que junto con el mensaje de error devuelve el valor de la MTU del enlace siguiente. Luego de recibir este mensaje el nodo de origen reduce el tamaño de los paquetes de acuerdo con la MTU indicada en el mensaje *packet too big*.

Este procedimiento finaliza cuando el tamaño del paquete es igual o menor que la MTU del camino, por lo que estas iteraciones (intercambio de mensajes y reducción del tamaño de los paquetes) pueden ocurrir varias veces hasta encontrar la menor MTU. Si el paquete es enviado a un grupo *multicast* el tamaño será la menor PMTU de todo el conjunto de destinos.

El PMTUD puede parecer imperfecto desde un punto de vista teórico, ya que el enrutamiento de los paquetes es dinámico y cada paquete puede ser entregado a través de una ruta diferente. Sin embargo, estos cambios no son tan frecuentes y si el valor de la MTU disminuye debido a un cambio de ruta el origen recibirá el mensaje de error y reducirá el valor de la *Path MTU*.

Más información:

- RFC 1981 - *Path MTU Discovery for IP version 6*

Jumbograms

- IPv6 permite enviar paquetes cuya longitud esté comprendida entre 65.536 y 4.294.967.295 bytes.
- Un *jumbogram* se identifica utilizando:
 - Valor 0 (cero) en el campo Tamaño de Datos.
 - *Hop-by-Hop* en el campo Siguiente Encabezado.
- El encabezado de extensión *Hop-by-Hop* indicará el tamaño de paquete.
- También se deben realizar modificaciones en los encabezados TCP y UDP, ambos limitados a 16 bits para indicar el tamaño máximo de los paquetes.

114

La RFC 2675 define una opción de encabezado de extensión *Hop-By-Hop* llamada Jumbo Payload. Esta opción permite enviar paquetes IPv6 con cargas útiles con una longitud comprendida entre 65.536 y 4.294.967.295 bytes, conocidos como *jumbograms*.

Al enviar *jumbograms*, el encabezado IPv6 indicará un valor nulo (cero) en los campos Tamaño de Datos y Siguiente Encabezado. Este último indicará que las opciones del encabezado de extensión *Hop-By-Hop* deben ser procesadas por los nodos, donde se indican los tamaños de los paquetes *jumbograms*.

El encabezado UDP tiene un campo de 16 bits llamado Tamaño que indica el tamaño del encabezado UDP más el tamaño de los datos y no permite el envío de paquetes de más de 65.536 bytes. No obstante, es posible enviar *jumbograms* definiendo el campo Tamaño como cero y permitiendo que el receptor determine el tamaño real del paquete UDP a partir del tamaño del paquete IPv6.

En los paquetes TCP las opciones *Maximum Segment Size* (MSS), que al iniciar la conexión negocia el tamaño máximo de paquete TCP a ser enviado, y *Urgent Pointer*, que indica el desplazamiento (offset) de bytes a partir del número de secuencia en que se encuentran los datos de alta prioridad, tampoco pueden referenciar paquetes mayores que 65.535 bytes. Así, para enviar *jumbograms* es necesario establecer el valor de MSS igual a 65.535, valor que el receptor del paquete tratará como infinito. La solución es similar para *Urgent Pointer*: se puede establecer un valor de *Urgent Pointer* igual a 65.535, indicando que apunta más allá del final de este paquete.

Más información:

- RFC 2675 - *IPv6 Jumbograms*

Gestión de Grupos Multicast

- MLD (*Multicast Listener Discovery*).
 - Equivalente a IGMPv2 en IPv4.
 - Utiliza mensajes ICMPv6.
 - Utiliza direcciones *link local* como dirección de origen.
 - Utiliza la opción *Router Alert* del encabezado de extensión *Hop-by-Hop*.
 - Nueva versión
 - MLDv2 (equivalente a IGMPv3).
- MRD (*Multicast Router Discovery*).
 - Mecanismo utilizado para descubrir routers *multicast*.
 - Utiliza 3 nuevos mensajes ICMPv6.

115

Recapitulando lo que se dijo en el módulo anterior, *multicast* es una técnica que permite direccionar múltiples nodos como un grupo, posibilitando el envío de paquetes a todos los nodos que lo componen a partir de una única dirección que lo identifica.

Los miembros de un grupo *multicast* son dinámicos y los nodos pueden entrar y salir de un grupo en cualquier momento. No existen limitaciones respecto del tamaño de un grupo *multicast*.

La gestión de los grupos *multicast* en IPv6 es realizada por el *Multicast Listener Discovery* (MLD), definido en la RFC 2710. Este protocolo es responsable de informar a los routers *multicast* locales el interés de los nodos en formar parte o salir de un determinado grupo *multicast*. En IPv4 este trabajo es realizado por el protocolo *Internet Group Management Protocol* (IGMPv2).

MLD utiliza tres tipos de mensajes ICMPv6:

- *Multicast Listener Query* (Tipo 130) – Hay dos subtipos de mensajes *Query*. Los mensajes *General Query* son utilizados por los routers para verificar periódicamente los miembros del grupo, solicitando a todos los nodos *multicast* que informen todos los grupos de los cuales forman parte. Los mensajes *Multicast-Address-Specific Query* son utilizados por los routers para descubrir si existen nodos que forman parte de determinado grupo;
- *Multicast Listener Report* (Tipo 131) – Un nodo envía mensajes *Report* no solicitados cuando éste comienza a formar parte de un grupo *multicast*. También son generados en respuesta a mensajes *Query*;
- *Multicast Listener Done* (Tipo 132) – Enviados por los nodos cuando abandonan determinado grupo.

Estos mensajes son enviados con una dirección de origen *link-local* y con valor 1 en el campo Límite de Direccionamiento, garantizando que permanezcan en la red local. Si el paquete tiene un encabezado *Hop-by-Hop*, el *flag Router Alert* estará marcado; de este modo los routers no descartarán el paquete aunque la dirección del grupo *multicast* en cuestión no esté siendo escuchada por los mismos.

En la RFC3810 se definió una nueva versión del protocolo MLD, denominada MLDv2. Equivalente a IGMPv3, además de incorporar las funcionalidades de gestión de grupos de MLD, esta nueva versión introduce el soporte para filtrado de origen, lo que permite que un nodo especifique si no desea recibir paquetes de un determinado origen o que informe su interés por recibir paquetes solo de determinadas direcciones. Por defecto, los miembros de un grupo reciben paquetes de todos los miembros de este grupo.

Otro mecanismo importante para el funcionamiento de los grupos *multicast* es el *Multicast Router Discovery* (MRD). Definido en la RFC 4286, el MRD es utilizado para descubrir routers *multicast* en la red. Utiliza tres mensajes ICMPv6:

- *Multicast Router Advertisement* (Tipo 151)– Este mensaje es enviado por los routers para anunciar que está habilitado el enrutamiento IP *multicast*. Se envía desde la dirección *link-local* del router a la dirección *multicast all-snoopers* (**FF02::6A**);
- *Multicast Router Solicitation* (Tipo 152) – Los dispositivos envían este mensaje a los routers *multicast* para solicitar mensajes *Multicast Router Advertisement*. Se envía desde la dirección *link-local* del dispositivo a la dirección *multicast all-routers* (**FF02::2**);
- *Multicast Router Termination* (Tipo 153) – Los routers envían este mensaje para anunciar que sus interfaces ya no están encaminando paquetes IP *multicast*. Se envía desde la dirección *link-local* del router a la dirección *multicast all-snoopers* (**FF02::6A**).

Todos los mensajes MRD también se envían con un Límite de Direccionamiento igual a 1 y con la opción *Router Alert*.

Más información:

- RFC 2710 - *Multicast Listener Discovery (MLD) for IPv6*
- RFC 4286 - *Multicast Router Discovery*

DNS

- Inmensa base de datos distribuida que se utiliza para resolver nombres de dominio en direcciones IP y viceversa.
- Arquitectura jerárquica en la cual los datos están dispuestos en forma de árbol invertido, distribuidos eficientemente en un sistema descentralizado y con caché.
- Registros
 - IPv4 = A - Traduce nombres a direcciones IPv4.
 - IPv6 = AAAA (quad-A) - Traduce nombres a direcciones IPv6.
- Ejemplo: www.ipv6.br. IN A **200.160.4.22**
 IN AAAA **2001:12ff:0:4::22**

117

El protocolo *Domain Name System* (DNS) es una inmensa base de datos distribuida que se utiliza para resolver nombres de dominio en direcciones IP y viceversa. Posee una arquitectura jerárquica en la cual los datos están dispuestos en forma de árbol invertido, distribuidos eficientemente en un sistema descentralizado y con caché.

En la RFC 3596 se definieron algunos cambios para permitir que el DNS trabaje con la versión 6 del protocolo IP.

Se creó un nuevo registro para almacenar las direcciones IPv6 de 128 bits, el AAAA o quad-A. Su función es traducir nombres a direcciones IPv6, equivalente al registro A que se utiliza en IPv4. Si un *host* tiene más de una dirección IPv6, éste tendrá un quad-A para cada dirección. Los registros se representan de la siguiente manera:

Ejemplo: www.ipv6.br. IN A **200.160.4.22**
 IN AAAA **2001:12ff:0:4::22**



Núcleo de Informação e Coordenação do Ponto BR




DNS

- Registro PTR – Resolución reversa.
 - IPv4 = in-addr.arpa - Traduce direcciones IPv4 a nombres.
 - IPv6 = ip6.arpa - Traduce direcciones IPv6 a nombres.

Ejemplo:

```

22.4.160.200.in-addr.arpa PTR www.ipv6.br.
2.2.0.0.0.0.0.0.0.0.0.0.0.0.4.0.0.0.0.0.0.f.f.2.1.1.0.0.2.ip6.arpa
PTR www.ipv6.br.
  
```

- Obsoletas
 - Registros
 - A6
 - DNAME
 - Dominio de resolución reversa
 - ip6.int

118



Para la resolución reversa se agregó el registro PTR ip6.arpa, responsable de traducir direcciones IPv6 a nombres. En su representación no se permite omitir la secuencia de ceros y el bit menos significativo se coloca más a la izquierda, como se puede observar en el siguiente ejemplo:

Exemplo:

```

22.4.160.200.in-addr.arpa PTR www.ipv6.br.
2.2.0.0.0.0.0.0.0.0.0.0.0.0.4.0.0.0.0.0.0.f.f.2.1.1.0.0.2.ip6.arpa
PTR www.ipv6.br.
  
```

Los otros tipos de registro DNS no sufrieron modificaciones, salvo que fueron adaptados para soportar el nuevo tamaño de las direcciones.

La RFC 2874 introdujo los registros A6 y DNAME a fin de facilitar la renumeración de redes, donde cada *nameserver* solamente tiene una parte de la dirección IPv6. Inicialmente el dominio de resolución reversa, definido en la RFC 1886, era ip6.int, pero hubieron manifestaciones contrarias a su utilización debido a que .int significa "internacional" y por lo tanto no debe ser empleado para fines administrativos en Internet. Los registros A6 y DNAME quedaron obsoletos por el desuso y el dominio .int fue reemplazado por el dominio .arpa, respectivamente en las RFC 3363 y 3152,

DNS

- La base de datos de un servidor DNS puede almacenar tanto registros IPv6 como IPv4.
- Estos datos son independientes de la versión de IP en que opera el servidor DNS.
 - Un servidor que solo tiene conexión IPv4 puede responder consultas AAAA o A.
 - Las informaciones obtenidas en la consulta IPv6 deben ser iguales a las obtenidas en la consulta IPv4.

119

Deben observarse dos aspectos del soporte para IPv6 del DNS. El primero es que un servidor DNS debe ser capaz de almacenar registros quad-A para direcciones IPv6. El segundo es que un servidor DNS es capaz de transportar consultas y respuestas a través de conexiones IPv6. Es decir, la base de datos de un servidor DNS puede almacenar tanto registros IPv6 como IPv4, independientemente de la versión de IP en que opera dicho servidor.

Por lo tanto, un servidor que solamente tiene conexión IPv4 puede responder tanto consultas AAAA como A. Sin embargo, las informaciones obtenidas en la consulta vía IPv6 deben ser iguales a las obtenidas en la consulta IPv4.

Más información:

RFC 3596 - *DNS Extensions to Support IP Version 6*

RFC 3363 - *Representing Internet Protocol version 6 (IPv6) Addresses in the Domain Name System (DNS)*

RFC 3364 - *Tradeoffs in Domain Name System (DNS) Support for Internet Protocol version 6 (IPv6)*

QoS

- El protocolo IP trata todos los paquetes de la misma manera, sin ningún tipo de preferencias.
- Algunas aplicaciones requieren que sus paquetes sean transportados con la seguridad de que habrá un retraso mínimo, latencia o pérdida de paquetes.
 - VoIP
 - Videoconferencias
 - Juegos en línea
 - Entre otros...
- Se utiliza el concepto de QoS (*Quality of Service*) o, en español, Calidad de Servicio.
- Principales arquitecturas: *Differentiated Services* (DiffServ) e *Integrated Services* (IntServ).
 - Ambas utilizan políticas de tráfico y se pueden combinar para permitir QoS en redes LAN o WAN.

Al principio, el protocolo IP trata todos los paquetes de la misma manera, sin ninguna preferencia a la hora de encaminarlos. Esto puede tener diferentes consecuencias para el desempeño de una aplicación, ya que en la actualidad muchas de estas aplicaciones tales como voz y video sobre IP, requieren transmisión y reproducción prácticamente en tiempo real y su calidad se puede reducir debido a la pérdida de paquetes, entrega fuera de orden, retraso o variación de la señal. Estos problemas pueden ocurrir debido a la forma en que el tráfico llega a los routers y es manipulado por los mismos, ya que como provienen de diferentes interfaces y redes el router procesa los paquetes en el orden en que se reciben.

El concepto de QoS (*Quality of Service*) o, en español, Calidad de Servicio, se utiliza para los protocolos cuya función es proporcionar la transmisión de determinados tráficos con prioridad y garantía de calidad. Actualmente existen dos arquitecturas principales: a *Differentiated Services* (DiffServ) e a *Integrated Services* (IntServ). Ambas utilizan políticas de tráfico y se pueden combinar para permitir QoS en redes LAN o WAN.

QoS

- DiffServ: Funciona por medio de clases, agregando y priorizando paquetes con requisitos de QoS similares.
 - IPv4 – Campo Tipo de Servicio (ToS).
 - IPv6 – Campo Clase de Tráfico:
 - Misma definición del campo ToS de IPv4.
 - Puede ser definido en el origen o por los routers.
 - Puede ser redefinido por los routers a lo largo del camino.
 - En los paquetes que no requieren QoS el campo Clase de Tráfico tiene valor 0 (cero).
- DiffServ no exige identificación ni gestión de flujos.
- Muy utilizado debido a su facilidad de implementación.

121

DiffServ trabaja por medio de clases, agregando y priorizando paquetes con requisitos de QoS similares.

Los paquetes *DiffServ* se identifican por los ocho bits de los campos Tipo de Servicio de IPv4 y Clase de Tráfico de IPv6, con el fin de identificar y distinguir las diferentes clases o prioridades de paquetes que requieren QoS.

Ambos campos tienen las mismas definiciones y las prioridades atribuidas a cada tipo de paquete se pueden definir tanto en el origen como en los routers, y también pueden ser redefinidas por los routers intermedios a lo largo del camino. En los paquetes que no requieren QoS el campo Clase de Tráfico tiene valor 0 (cero).

En comparación con *IntServ*, *DiffServ* no exige ninguna identificación ni gestión de los flujos y en general es más utilizado en las redes debido a su facilidad de implementación.

QoS

- IntServ: Se basa en la reserva de recursos por flujo. Normalmente está asociado al RSVP (*Resource ReSerVation Protocol*).
- IPv6 – El campo Identificador de Flujo es completado por el origen con valores aleatorios comprendidos entre 00001 y FFFFF para identificar el flujo que requiere QoS.
 - Los paquetes que no pertenecen a un flujo deben tener este campo completado con ceros.
 - Los *hosts* y routers que no tienen soporte para las funciones del campo Identificador de Flujo deben completar este campo con ceros cuando envían un paquete, no modificarlo cuando encaminan un paquete o ignorarlo cuando reciben un paquete.
- Los paquetes de un mismo flujo deben tener la misma dirección de origen y destino y el mismo valor en el campo Identificador de Flujo.
- RSVP utiliza algunos elementos del protocolo IPv6, como por ejemplo el campo Identificador de Flujo y el encabezado de extensión *Hop-by-Hop*.

El modelo IntServ se basa en la reserva de recursos por flujo y su utilización normalmente está asociada al protocolo RSVP. El protocolo RSVP se utiliza para reservar el recurso a lo largo del camino de un flujo que requiere QoS, desde la fuente hasta el destino.

En IPv6, para identificar los flujos que requieren QoS se utilizan los 20 bits del campo Identificador de Flujo, que se completan con valores aleatorios comprendidos entre 00001 y FFFFF. Los paquetes que no pertenecen a un flujo deben completar el campo Identificador de Flujo con ceros. Los *hosts* y routers que no tienen soporte para las funciones de este campo deben completarlo con ceros cuando envían un paquete, no modificarlo cuando encaminan un paquete o ignorarlo cuando reciben un paquete. Los paquetes de un mismo flujo deben tener la misma dirección de origen y destino y el mismo valor en el campo Identificador de Flujo.

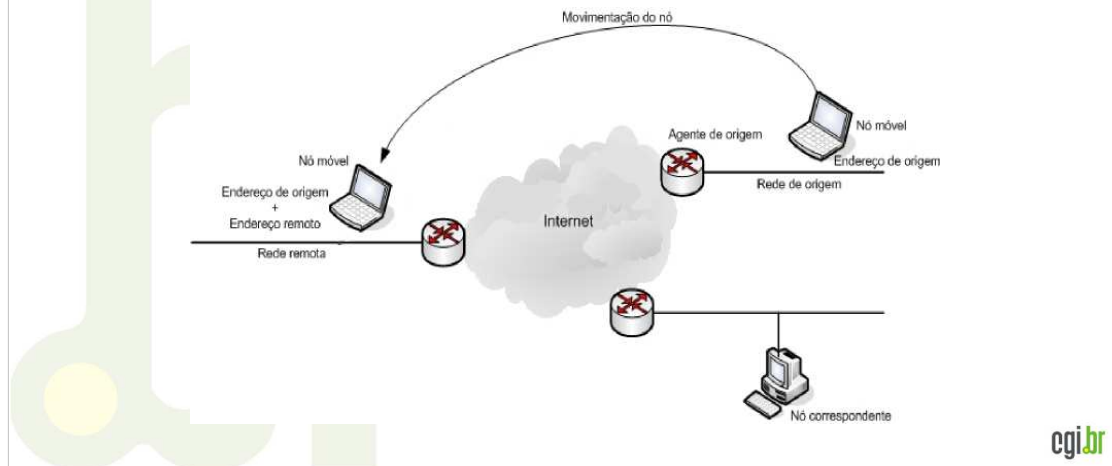
RSVP utiliza algunos elementos del protocolo IPv6, como por ejemplo el campo Identificador de Flujo y el encabezado de extensión *Hop-by-Hop*. Los paquetes RSVP se envían con el mismo valor en el campo Identificador de Flujo, junto con el encabezado de extensión *Hop-By-Hop*, usado para transportar un mensaje *Router Alert* que le indica a cada router en el camino del tráfico QoS que el paquete IP debe ser procesado.

Más información:

- RFC 1633 - *Integrated Services in the Internet Architecture: an Overview*
- RFC 2205 - *Resource ReSerVation Protocol (RSVP)*
- RFC 2475 - *An Architecture for Differentiated Services*
- RFC 3260 - *New Terminology and Clarifications for Diffserv*

Movilidad IPv6

- Permite que un dispositivo móvil se desplace de una red a otra sin necesidad de cambiar su dirección IP de origen, haciendo que el movimiento entre redes sea invisible para los protocolos de las capas superiores.



El soporte para movilidad permite que un dispositivo móvil se desplace de una red a otra sin necesidad de cambiar su dirección IP de origen, haciendo que el movimiento entre redes sea invisible para los protocolos de las capas superiores. Por lo tanto, todos los paquetes enviados a este nodo móvil continuarán siendo encaminados al mismo usando la dirección de origen.

Existen algunos componentes clave para el funcionamiento del soporte para movilidad IPv6:

- **Nodo móvil** – Dispositivo que puede cambiar de una red a otra sin dejar de recibir paquetes a través de su Dirección de Origen;
- **Red de origen** – Red que atribuye la Dirección de Origen al Nodo Móvil;
- **Agente de Origen** – Router ubicado en la Red de Origen y que mantiene la asociación entre la Dirección de Origen y la Dirección Remota del Nodo Móvil.
- **Dirección de origen** – Dirección *global unicast* atribuida por la Red de Origen al Nodo Móvil. Se utiliza como dirección permanente hacia la cual se encaminan los paquetes.
- **Red remota** – Cualquier red diferente a la de Red de Origen en la cual se encuentra el Nodo Móvil;
- **Dirección remota** – Dirección *global unicast* atribuida al Nodo Móvil por la Red remota;
- **Nodo correspondiente** – Nodo que se comunica con un Nodo Móvil. Puede ser móvil o estacionario.

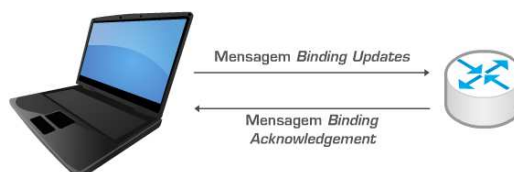
Movilidad IPv6

- Funcionamiento

- El Nodo Móvil utiliza la Dirección de Origen para recibir paquetes en la Red de Origen.

- Desplazamiento

- Adquiere Dirección Remota mediante autoconfiguración *stateless* o *stateful*.
- El Agente de Origen realiza la asociación entre la Dirección Remota y la Dirección de Origen.



- El Nodo Móvil se puede registrar directamente con el Nodo Correspondiente.

124

El Nodo Móvil tiene una Dirección de Origen fija, que le es atribuida por su Red de Origen. Esta dirección se mantiene aunque no se desplace de su Red de Origen.

Al ingresar en una Red Remota el Nodo Móvil recibe una o más Direcciones Remotas a través de los mecanismos de autoconfiguración, que consisten en un prefijo válido en la Red Remota. Para asegurar que se reciban los paquetes IPv6 destinados a su Dirección de Origen, el nodo realiza una asociación entre la Dirección de Origen y la Dirección Remota, registrando su nueva dirección en el Agente de Origen mediante el envío de un mensaje *Binding Updates*. Como respuesta a este mensaje el router de la Red de Origen envía un mensaje *Binding Acknowledgement*.

Esta asociación de direcciones también se puede realizar directamente con el Nodo Correspondiente a fin de optimizar la comunicación.

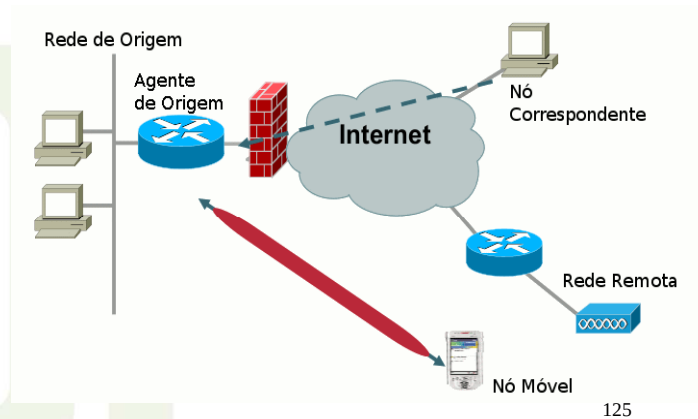
Para detectar que regresó a su red el Nodo Móvil utiliza el proceso de Descubrimiento de Vecinos Inaccesibles para determinar si su router por defecto está activo. En caso de ubicar un nuevo router por defecto, éste generará una nueva dirección basada en el prefijo anunciado en el mensaje RA. Sin embargo, el hecho de encontrar un nuevo router por defecto no necesariamente significa que se encuentra en una nueva red; puede que simplemente se trate de una renumeración de la red o del agregado de un nuevo router. Por lo tanto, antes de asociar las direcciones con el Agente de Origen y con los Nodos Correspondientes, el Nodo Móvil intentará ubicar nuevamente su router por defecto y comparará si el intervalo entre el envío de mensajes RA no solicitados es el mismo que está configurado en su Red Original.

Cuando el Nodo Móvil regresa a su Red de Origen envía un mensaje *Binding Updates* para informarle su retorno al Agente de Origen y avisarle que ya no es necesario que encamine los paquetes.

Movilidad IPv6

- El encaminamiento de los paquetes hacia el Nodo Móvil puede producirse de dos maneras diferentes:

- **Túneles bidireccionales**



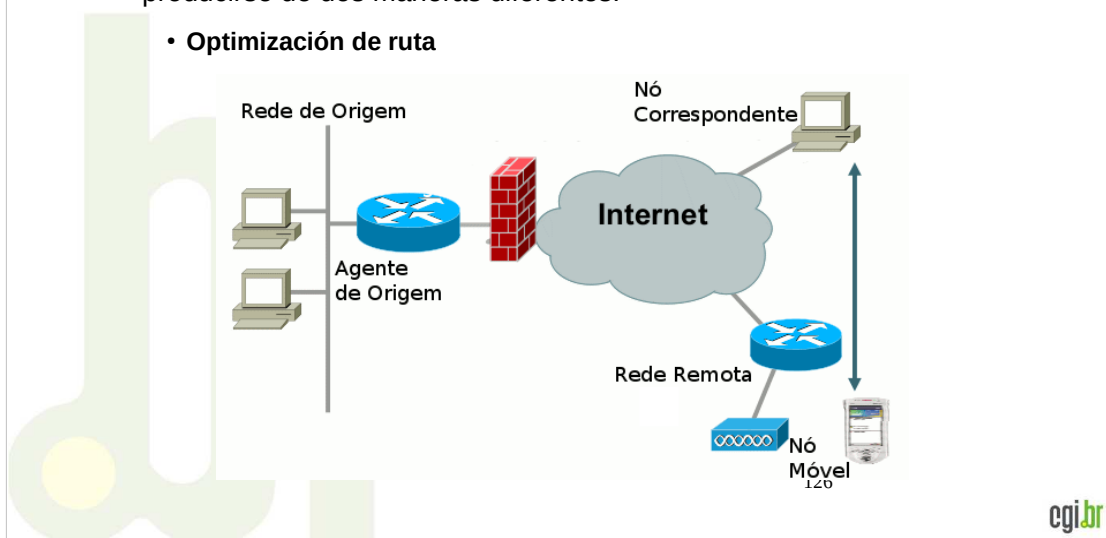
Las comunicaciones entre los nodos móviles y los nodos correspondientes pueden producirse de dos maneras diferentes: túneles bidireccionales y optimización de ruta.

En el caso de los túneles bidireccionales, los paquetes que el Nodo Correspondiente envía a la Dirección Original del Nodo Móvil son interceptados por el Agente de Origen, que los encamina, a través de un túnel, hacia el Nodo Móvil utilizando la Dirección Remota. Luego el Nodo Móvil responde al Agente de Origen, a través del túnel, que reenvía el paquete al Nodo Correspondiente. En este caso no es necesario que el Nodo Correspondiente tenga soporte para movilidad IPv6 ni que el Nodo Móvil esté registrado en el Nodo Correspondiente.

Movilidad IPv6

- El encaminamiento de los paquetes hacia el Nodo Móvil puede producirse de dos maneras diferentes:

- **Optimización de ruta**



En el modo Optimización de Ruta la comunicación entre el Nodo Móvil y el Nodo Correspondiente ocurre de manera directa, sin necesidad de utilizar el Agente de Origen. Para que esta comunicación ocurra el Nodo Móvil registra su Dirección Remota en el Nodo Correspondiente, que asocia las Direcciones de Origen y Remota del Nodo Móvil.

El intercambio de mensajes entre ambos nodos funciona de la siguiente manera:

- El Nodo Correspondiente envía paquetes con el campo Dirección de Destino del encabezado base completado con la Dirección Remota del Nodo Móvil. El encabezado base es seguido por el encabezado de extensión *Routing* Tipo 2, que transporta la Dirección de Origen del Nodo Móvil;
- Al recibir el paquete el Nodo Móvil procesa el encabezado *Routing* e inserta la Dirección de Origen del encabezado *Routing* en el campo Dirección de Destino del encabezado base. Las capas superiores continúan procesando el paquete normalmente;
- Los paquetes enviados por el Nodo Móvil tiene el campo Dirección de Origen del encabezado base completado con la Dirección Remota. El encabezado base está seguido por el encabezado de extensión *Destination Options*, que transporta en la opción *Home Address* la Dirección de Origen del Nodo Móvil;
- Al recibir el paquete el Nodo Correspondiente inserta la Dirección de Origen del encabezado *Destination Options* en el campo Dirección de Origen del encabezado base. Las capas superiores continúan procesando el paquete normalmente.

Movilidad IPv6

- Identificado por el valor 135 en el campo Siguiente Encabezado.
- Se utiliza en el intercambio de mensajes relacionados con la creación y gestión de las asociaciones de direcciones.
- Encabezado de extensión *Mobility*

Protocolo de datos	Tam. encab. de extensión	Tipo de mensaje <i>Mobility</i>	Reservado
Suma de verificación		Datos	

- Principales tipos de mensajes *Mobility*:
 - *Binding Refresh Request* (Tipo 0)
 - *Binding Update* (Tipo 5)
 - *Binding Ack* (Tipo 6)
 - *Binding Error* (Tipo 7)

127

Para optimizar el funcionamiento de este servicio, a la especificación de IPv6 se agregó un nuevo encabezado de extensión, el encabezado *Mobility*.

El encabezado de extensión *Mobility* se identifica por el valor 135 en el campo Siguiente Encabezado. Es utilizado por el Nodo Móvil, por el Agente Remoto y por el Nodo Correspondiente en el intercambio de mensajes relacionados con la creación y gestión de asociaciones de direcciones.

Este encabezado tiene los siguientes campos:

- Protocolo de datos – Corresponde al campo Siguiente Encabezado. En la actualidad solo se utiliza el valor 59, en formato decimal, lo que indica que no hay encabezado siguiente;
- Tamaño del encabezado de extensión – Contiene el tamaño del encabezado *Mobility* en unidades de 8 bytes. El tamaño de este encabezado debe ser siempre múltiplo de 8;
- Tipo de Mensaje *Mobility* – Indica el tipo de mensaje enviado;
- Suma de Verificación – Verifica la integridad del encabezado *Mobility*;
- Datos – Su formato y tamaño dependen del tipo de mensaje *Mobility* que está siendo enviado.

Los siguientes son los tipos de mensajes *Mobility* más utilizados:

- *Binding Refresh Request* (Tipo 0) – Enviado por el Nodo Correspondiente para solicitarle al Nodo Móvil la actualización de la asociación de direcciones;
- *Binding Update* (Tipo 5) – Enviado por el Nodo Móvil para notificar una nueva Dirección Remota al Agente de Origen o al Nodo Correspondiente;
- *Binding Ack* (Tipo 6) – Enviado para confirmar la recepción de un mensaje *Binding Update*;
- *Binding Error* (Tipo 7) – Enviado por el Nodo Correspondiente para informar errores.

Movilidad IPv6

- Nuevos mensajes ICMPv6
 - *Home Agent Address Discovery Request*;
 - *Home Agent Address Discovery Reply*;
 - *Mobile Prefix Solicitation*;
 - *Mobile Prefix Advertisement*.

128

También se crearon cuatro nuevos mensajes ICMPv6 que se utilizan en la configuración de prefijos en la Red de Origen y en el descubrimiento de Agentes de Origen.

El par de mensajes *Home Agent Address Discovery Request* y *Home Agent Address Discovery Reply* se utilizan para descubrir dinámicamente un Agente de Origen en su red. Esto evita la necesidad de realizar configuraciones manuales y problemas en caso de reenumeración del Agente de Origen.

El Nodo Móvil envía un mensaje *Discovery Request* a la dirección *anycast* del Agente de Origen en su red. El campo Dirección de Origen del encabezado base lleva la Dirección Remota del Nodo Móvil. El Agente de Origen responde enviando un mensaje *Discovery Reply*. Los mensajes *Discovery Request* y *Reply* se identifican, respectivamente, por los valores 150 y 151 en el campo Siguiente Encabezado.

Ya el mensaje *Mobile Prefix Solicitation* es enviado por el Nodo Móvil al Agente de Origen para determinar cambios en la configuración del prefijo en la red. El Agente Remoto responde enviando un mensaje *Mobile Prefix Advertisement*. En base a esta respuesta el Nodo Móvil puede configurar su Dirección de Origen. Los mensajes *Mobile Prefix Solicitation* y *Advertisement* se identifican, respectivamente, en el campo Siguiente Encabezado por los valores 152 y 153.

Movilidad IPv6

- Cambios en el Protocolo de Descubrimiento de Vecinos:

- Modificación del formato de los mensajes RA;
- Modificación del formato de *Prefix Information*;
- Agregado de la opción *Advertisement Interval*;
- Agregado de la opción *Home Agent Information*.

129

También se introdujeron algunas modificaciones en el Protocolo de Descubrimiento de Vecinos.

A los mensajes RA se agregó el *flag* H, que permite que un router anuncie si actúa como Agente de Origen en la red. Basándose en este anuncio, un Nodo Móvil crea una lista de Agentes de Origen de su red.

Así y todo, para mantener esta lista actualizada el Nodo Móvil debe conocer las direcciones *global unicast* de los routers, pero los mensajes RA traen solamente la dirección *link-local*. Para resolver este problema se agregó un nuevo *flag* a la opción *Prefix Information*, el *flag* R. Cuando este *flag* está marcada indica que la opción *Prefix Information* no contiene un prefijo pero sí la dirección *global unicast* del router.

También se crearon dos nuevas opciones para el protocolo, *Advertisement Interval* y *Home Agent Information*. La primera indica el intervalo entre mensajes RA no solicitados, información que es utilizada por el algoritmo de detección de cambio de red. De acuerdo con las especificaciones del Protocolo de Descubrimiento de Vecinos el intervalo mínimo entre el envío de estos mensajes debe ser de tres segundos. Sin embargo, para asegurar que el Nodo Móvil detecte el cambio de red y aprenda la información sobre la nueva red lo más rápidamente posible, los routers con soporte para movilidad IPv6 se pueden configurar con un intervalo de tiempo menor para el anuncio de mensajes RA.

La opción *Home Agent Information* se utiliza para indicar el nivel de preferencia de asociación de cada Agente de Origen.

Movilidad IPv6

- Movilidad IPv4 x Movilidad IPv6:
 - No requiere la implementación de Agentes Remotos;
 - La optimización de la ruta está incorporada en el protocolo;
 - La autoconfiguración *stateless* facilita la atribución de Direcciones Remotas;
 - Aprovechas los beneficios del protocolo IPv6:
 - Descubrimiento de Vecinos, ICMPv6, encabezados de extensión...
 - Utiliza el protocolo de Descubrimiento de Vecinos en lugar de ARP;
 - Utiliza *anycast* para localizar Agentes de Origen en lugar de *broadcast*.

130

Las principales diferencias entre el soporte para movilidad de IPv6 y de IPv4 se pueden resumir en los siguientes puntos:

- Ya no es necesario implementar routers especiales que actúen como agentes remotos;
- En lugar de formar parte de un conjunto de extensiones opcionales, la optimización de la ruta está incorporada en el protocolo;
- La autoconfiguración *stateless* facilita la atribución de Direcciones Remotas;
- Aprovecha los beneficios del protocolo IPv6 tales como el protocolo de Descubrimiento de Vecinos, los mensajes ICMPv6 y los encabezados de extensión.
- El uso del protocolo de Descubrimiento de Vecinos en lugar de ARP permite que el proceso de intercepción de los paquetes destinados al nodo móvil no dependa de la capa de enlace, lo que simplifica el protocolo y aumenta su eficiencia;
- La búsqueda por agentes de origen que realizaba el nodo móvil pasó a ser realizada utilizando *anycast*. De esta forma el nodo móvil solo recibirá la respuesta de un único agente de origen. Con IPv4 se utiliza *broadcast*, lo que implica una respuesta separada para cada agente de origen existente.

Más información:

- RFC 3775 - *Mobility Support in IPv6*

IPv6.br

La nueva generación del
Protocolo de Internet

Administración y Monitoreo de Redes IPv6

Módulo 5

132

egi.br

El uso de protocolos y herramientas de administración y monitoreo de redes es muy importante para mantener la calidad y lograr el máximo desempeño de la red. La adopción del nuevo Protocolo de Internet requiere saber cuáles de estas herramientas pueden recolectar información acerca de IPv6 y si se pueden obtener a través de la red IPv6.

En este módulo estos aspectos se abordarán en relación con algunos protocolos utilizados para estos fines, tales como SSH, FTP, SNMP, entre otros, y algunas aplicaciones como Argus, Nagios, MRTG, Rancid, Wireshark y Looking Glass.

Administración y Monitoreo

- Necesario para mantener la calidad de la red.
- Se realiza a través de diferentes herramientas y protocolos.
 - Deben cubrir diferentes segmentos:
 - LAN
 - WAN
 - Abarcar diferentes aspectos:
 - Acceso remoto
 - Información sobre el flujo de datos
 - Seguridad
 - Mantenimiento
 - Acceso a la información
- Deben recolectar información sobre IPv6 y transmitirla vía conexiones IPv6.

133

Requisito fundamental para mantener la calidad y garantizar su máxima eficiencia, la administración y monitoreo de redes de computadoras es una parte importante de su operación sin importar cuál sea el tamaño de la red.

Actualmente existen incontables herramientas y protocolos que realizan estas funciones, los cuales se diferencian de acuerdo con el segmento de red en el que actúan (LAN o WAN) y las funcionalidades que ofrecen, como por ejemplo garantizar el acceso remoto y seguro a los nodos de la red, recolectar información acerca del flujo de datos, autenticación de usuarios, pruebas y mantenimiento, y acceso a la información.

Por lo tanto, en este momento de transición IPv4 e IPv6, es importante que estas herramientas sean capaces de soportar ambas versiones del protocolo IP y puedan recolectar información sobre IPv6 y transmitirla vía conexiones IPv6.

Funciones básicas

- Funciones básicas de la administración de redes
 - Acceso remoto:
 - SSH;
 - TELNET.
 - Transferencia de archivos
 - SCP;
 - FTP;
 - TFTP.

134

Una de las funciones más básicas de la administración de redes es el acceso remoto a otros dispositivos. En este sentido, los principales protocolos existentes ya son capaces de operar sobre IPv6.

Los protocolos Telnet y SSH (Secure Shell), utilizados para establecer conexiones remotas con otros dispositivos de la red, ya permiten el acceso vía conexiones IPv6. Aplicaciones como OpenSSH y PuTTY, por ejemplo, ya ofrecen esta funcionalidad.

Del mismo modo, ya es posible transferir archivos entre dispositivos remotos vía IPv6 a través de protocolos como SCP, TFTP y FTP. FTP incluso fue uno de los primeros protocolos en ser adaptados para trabajar sobre IPv6.

SNMP y las MIB

- SNMP: protocolo más utilizado en la administración de redes IPv4.
- Su funcionamiento se basa en la utilización de dos dispositivos: agentes y administradores.
- El administrador obtiene la información enviando solicitudes a uno o más agentes.
- La información puede ser transportada tanto mediante conexiones IPv4 como mediante conexiones IPv6.
 - El tipo de información transportada (IPv4 o IPv6) es independiente del protocolo de red utilizado en la conexión.
 - Desde 2002 han existido implementaciones sobre IPv6.
- MIB: estructura de datos que modela toda la información necesaria para la administración de la red.

135

En las redes IPv4, el protocolo SNMP (*Simple Network Management Protocol*), definido en la RFC 1157, es una de las herramientas de administración más utilizadas por su flexibilidad y facilidad de implementación.

El funcionamiento de SNMP se basa en la utilización de dos dispositivos, un agente y un administrador. Cada dispositivo administrado debe tener un agente y una base de datos referente a su estado actual que puede ser consultada y modificada por el administrador. El conjunto de estos datos, u objetos administrados, se conoce como MIB (*Management Information Base*), una estructura de datos que modela toda la información necesaria para la administración de la red.

El agente es responsable por el mantenimiento de la información administrada. Es quien debe responder a las solicitudes del administrador, enviándole la información necesaria para que éste pueda realizar el monitoreo del sistema.

El envío de esta información, almacenadas en las MIB, se puede realizar tanto mediante conexiones IPv4 o IPv6, ya que el tipo de información transportada es independiente del protocolo de red utilizado. Sin embargo, desde el año 2002 existen implementaciones de SNMP capaces de monitorear redes que solo poseen conexiones IPv6.

SNMP y las MIB

- Es necesario que las MIB puedan reconocer información sobre la red IPv6.
- 1998: Se define un enfoque solo para direcciones IPv6. Pero era necesario implementar una MIB para cada protocolo.
- 2006: Se elaboró una MIB unificada, creando un único conjunto de objetos capaz de describir y administrar módulos IP de forma independiente del protocolo.
 - *InetAddressType*
 - *InetAddress*

136

Aunque el tipo de protocolo de red utilizado en la transmisión de datos no interfiere con el envío de mensajes SNMP, es necesario que las MIB sean capaces de almacenar información sobre la red IPv6.

Por este motivo, en 1998, en la RFC 2465 (que se volvió obsoleta con la RFC 4292) se definió un enfoque solo para direcciones IPv6. Pero era necesario implementar una MIB para cada protocolo. En 2002, la RFC 2851 (que se volvió obsoleta con las RFC 3291 y 4001) estableció una MIB unificada, creando así un único conjunto de objetos capaz de describir y administrar módulos IP de forma independiente del protocolo.

Esta nueva convención una dirección IP como una estructura *{inetAddressType, inetAddress}*, donde el primero es un valor entero que determina la forma en que será codificado el segundo.

Más información:

- RFC 1157 - *A Simple Network Management Protocol (SNMP)*
- RFC 4001 - *Textual Conventions for Internet Network Addresses*
- RFC 4292 - *IP Forwarding Table MIB*

Monitoreo de Flujo

- Flujo - Conjunto de paquetes pertenecientes a la misma aplicación que tienen la misma dirección de origen y de destino.
- Los equipos de red envían información sobre un determinado flujo de datos hacia el colector, que almacena e interpreta dichos datos.
- NetFlow - Protocolo desarrollado por Cisco Systems, ya incluye soporte para IPv6.
- IPFIX - Basado en el protocolo NetFlow, también es capaz de exportar y recolectar datos sobre el tráfico IPv6.

137

cgi.br

Para realizar análisis más detallados de una red podemos aplicar un enfoque alternativo que consiste en recolectar información acerca de cada paquete. En este método, los equipos de red, por ejemplo un router, envían periódicamente información sobre un determinado flujo de datos a un dispositivo llamado colector que almacena e interpreta estos datos.

Un flujo de datos se puede definir como un conjunto de paquetes pertenecientes a la misma aplicación que tienen la misma dirección de origen y destino. Los principales protocolos utilizados para transmitir información sobre un flujo IP de una red también están preparados para recolectar datos sobre el tráfico IPv6.

El protocolo NetFlow, desarrollado por Cisco Systems y definido en la RFC 3954, es una herramienta eficiente que se utiliza, entre otras cosas, para contabilizar y caracterizar el tráfico de las redes, planear redes y detectar ataques DoS y DDoS. El protocolo NetFlow con soporte para IPv6 está implementado a partir de Cisco IOS 12.3(7)T, pero esta implementación todavía utiliza el protocolo IPv4 para la exportación de datos.

Del mismo modo, el protocolo IPFIX (*IP Flow Information Export*) propuesto por la IETF en la RFC 3917 también puede exportar y recolectar datos sobre el tráfico IPv6.

Más información:

- RFC 3917 - *Requirements for IP Flow Information Export* (IPFIX)
- RFC 3954 - Cisco Systems NetFlow Services Export V9

NTP

- La sincronización de los relojes de los equipos puede repercutir significativamente en el funcionamiento de las redes.
- El protocolo NTP mantiene el reloj del equipo siempre en la hora correcta, con una exactitud de algunas milésimas de segundo.
- Esto se puede hacer sincronizando el reloj del equipo con un servidor NTP público.
- Servidores NTP públicos en Brasil con soporte para IPv6
 - a.ntp.br
 - ntp.pop-sc.rnp.br
 - ntp.pop-rs.rnp.br
 - ntp.cert-rs.tche.br
 - ntp.pop-mg.rnp.br

138

Un elemento muy importante en la gestión de redes, la sincronización de los relojes de los equipos puede repercutir de significativamente en el funcionamiento de diferentes programas y sistemas, como así también en la seguridad de los equipos, las redes y la propia Internet. La utilización del protocolo NTP (*Network Time Protocol*) permite mantener el reloj del equipo siempre en la hora correcta, con una exactitud de algunas milésimas de segundo. Esto se puede hacer sincronizando el reloj del equipo con un servidor NTP público.

Los relojes de una red IPv6 se pueden sincronizar conectándose a servidores que posean soporte para el nuevo protocolo IP. La siguiente lista muestra las direcciones de una serie de servidores NTP públicos en Brasil que ya trabajan con conexión IPv6:

- a.ntp.br
- ntp.pop-sc.rnp.br
- ntp.pop-rs.rnp.br
- ntp.cert-rs.tche.br
- ntp.pop-mg.rnp.br

Más información:

- RFC 1305 - *Network Time Protocol (Version 3) Specification, Implementation and Analysis*

Herramientas de Monitoreo

- ARGUS
 - Soporte para IPv6 a partir de la versión 3.2;
 - Aplicación para monitoreo de redes y sistemas
 - Permite recolectar y evaluar datos sobre:
 - Conectividad en la red;
 - Puertos TCP/UDP;
 - Aplicaciones - HTTP, SMTP, RADIUS, etc.
- NAGIOS
 - Herramienta versátil y flexible;
 - Principales funcionalidades:
 - Monitoreo de servicios de red;
 - Monitoreo de recursos de los *hosts*;
 - Notificación de errores;
 - Adición de nuevas funcionalidades a través de *plugins*;
 - Soporte para IPv6 incluido en las versiones de *plugins* 1.4.x.

139

Existen diferentes herramientas que ayudan a monitorear una red. Utilidades para gestión de tráfico, elaboración de gráficos e informes sobre el estado de los equipos y enlaces, y análisis de tráfico son algunas de las herramientas utilizadas frecuentemente por los administradores de redes, y muchas de ellas ya tienen soporte para IPv6. Entre estas herramientas podemos destacar:

- ARGUS – Aplicación para monitoreo de redes y sistemas que permite recolectar y evaluar datos relacionados con la conectividad en la red, puertos TCP/UDP y aplicaciones como HTTP, SMTP, RADIUS, etc. Incluye soporte para IPv6 a partir de la versión 3.2;
- NAGIOS – Herramienta versátil y flexible que tiene múltiples funcionalidades tales como monitoreo de servicios de red; de recursos de los *hosts*; notificación de errores; etc.. Tiene la ventaja de permitir la adición de nuevas funcionalidades a través de *plugins*. Las versiones 1.4.x de los *plugins* incluyen soporte para IPv6;

Herramientas de Monitoreo

- NTOP
 - Detallar la utilización de la red;
 - Visualización de estadísticas del tráfico;
 - Análisis del tráfico IP;
 - Detección de violaciones de seguridad;
 - Tiene soporte para tráfico IPv6.
- MRTG
 - Desarrollado en C y Perl;
 - Utiliza SNMP para obtener información de los dispositivos administrados;
 - Análisis de los datos mediante gráficos visualizados en formato HTML;
 - Soporte para IPv6 a partir de la versión 2.10.0.

140

- NTOP (*Network Traffic Probe*) – Puede detallar la utilización de la red por *host*, protocolo, etc., permitiendo visualizar estadísticas de tráfico, análisis del tráfico IP, detección de violaciones de seguridad en la red, entre otras funciones. Tiene soporte para tráfico IPv6;
- MRTG (*Multi Router Traffic Grapher*) – Desarrollado en C y Perl, utiliza SNMP para obtener información del tráfico de los dispositivos administrados. Todos los datos obtenidos a través del protocolo SNMP pueden ser monitoreados por esta herramienta y analizados a través de gráficos visualizados en formato HTML. Soporte para IPv6 a partir de la versión 2.10.0.

Herramientas de Monitoreo

- Pchar
 - Herramienta utilizada para evaluar rendimiento;
 - Análisis de ancho de banda;
 - Análisis de latencia;
 - Análisis de pérdida de conexiones;
 - Permite analizar redes IPv6.
- Rancid
 - Monitorea la configuración de los equipos;
 - Desarrollada en lenguajes Perl, Shell y C;
 - Proporciona un *looking glass*;
 - Puede caracterizar el camino entre dos *hosts* en redes IPv6.

141

- Pchar – Herramienta para evaluar el rendimiento de la red. Analiza aspectos tales como el ancho de banda, la latencia y la pérdida de conexiones. Permite analizar redes IPv6.
- Rancid – Permite monitorear la configuración de los equipos (*software* y *hardware*) utilizando CVS. Desarrollada en lenguajes Perl, Shell y C, además de las funcionalidades tradicionales de las herramientas para monitoreo de redes, Rancid proporciona un *looking glass*. Puede caracterizar el camino entre dos *hosts* en redes IPv6.

Herramientas de Monitoreo

- Wireshark
 - Analizador de tráfico de red (*sniffer*);
 - Posee interfaz gráfica;
 - Presenta información sobre:
 - Árbol de protocolos de los paquetes;
 - Contenido de los paquetes;
 - Permite la captura de paquetes IPv6.
- Looking Glass
 - Permite obtener información sobre los routers sin necesidad de acceder directamente a los equipos;
 - Se puede acceder a esta herramienta a través de una interfaz web, lo que facilita el diagnóstico de problemas en la red;
 - Permite el acceso mediante conexiones IPv6.

142

- Wireshark – Analizador de tráfico de red (*sniffer*) a través de la captura de paquetes. Posee una interfaz gráfica y presenta información sobre el árbol de protocolos de los paquetes y su contenido. Permite la captura de paquetes IPv6;
- Looking Glass – Permite obtener información sobre un router sin necesidad de acceder directamente al equipo; Se puede acceder a esta herramienta a través de una interfaz web, lo que facilita el diagnóstico de problemas en la red. Permite el acceso mediante conexiones IPv6.

IPv6.br

La nueva generación del
Protocolo de Internet

143

egi.br

Seguridad

Módulo 6

144

En el proyecto de IPv6 el tema de la seguridad se consideró desde el inicio. Los mecanismos de autenticación y encriptación pasaron a formar parte del protocolo IPv6, poniendo a disposición de cualquier par de dispositivos de una conexión end-to-end métodos que buscan garantizar la seguridad de los datos que atraviesan la red. Sin embargo, todavía existen muchos problemas y también han surgido nuevas fallas de seguridad.

En este módulo abordaremos cada uno de estos nuevos escenarios, analizando las nuevas herramientas de seguridad de IPv6 y trazando un paralelo entre las amenazas que existen en IPv6 y en su predecesor.

Seguridad

- IPv4
 - Pensado para interconectar redes académicas – sin mucha preocupación por la seguridad.
 - Uso comercial – operaciones bancarias, comercio electrónico, intercambio de información confidencial.....
 - Amenazas
 - Barrido de direcciones (*scanning*)
 - Falsificación de direcciones (*spoofing*)
 - Manipulación de encabezados y fragmentación
 - Virus, troyanos y gusanos
 - ...
 - NAT + IPSec son incompatibles

145

egi.br

Destinado principalmente a interconectar redes de investigación académicas, el proyecto original de IPv4 no presentaba ninguna gran preocupación con respecto de la seguridad de la información transmitida. Sin embargo, la creciente importancia de Internet para las transacciones entre empresas y consumidores llevó a exigir mayores niveles de seguridad, como por ejemplo identificación de usuarios y encriptación de los datos, haciendo que fuera necesario agregar al protocolo original nuevos mecanismos que garantizaran estos servicios. No obstante, las soluciones adoptadas son habitualmente específicas para cada aplicación.

Este hecho es bastante claro en la Internet actual. Hay quienes sostienen que para resolver este problema debería haber una nueva Internet, proyectadas desde cero (enfoque *clean slate*).

Seguridad en IPv6

- ¿IPv6 es más seguro?
 - Presenta nuevos problemas:
 - Técnicas de transición;
 - Descubrimiento de vecinos y autoconfiguración;
 - Modelo *end-to-end*;
 - Movilidad IPv6;
 - Falta de “*Best Practices*”, políticas, entrenamiento, herramientas....

146

Al proyectar IPv6 se abordaron algunos aspectos de la seguridad, pero sus implementaciones aun no están maduras. A pesar de tener más de diez años, no hay buena experiencia en su uso. Las mejores prácticas continúan siendo tomadas de IPv4, por lo que no siempre funcionan bien.

Seguridad en IPv6

- ¿IPv6 es más seguro?
 - Herramientas de seguridad
 - IPSec
 - *Secure Neighbor Discovery* (SEND)
 - Estructura de las direcciones
 - *Cryptographically Generated Address* (CGA)
 - Extensiones de privacidad
 - *Unique Local Addresses* (ULA)

147

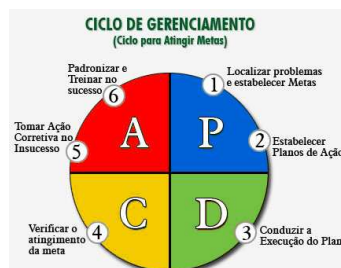
En IPv6 la seguridad fue una preocupación desde el inicio, por lo que en el protocolo se implementaron varias herramientas de seguridad.

Más información:

- RFC 4864 - *Local Network Protection for IPv6*

Seguridad en IPv6

- Estrategia de implementación
 - Sin planificación... Ejemplo:
 - Routers inalámbricos
 - Sistemas sin *firewall*
 - Proyectos de última hora / demostraciones
 - Proyectos sin la participación de especialistas en seguridad
 - O...
 - *Planificar (Plan)*
 - *Implementar (Do)*
 - *Verificar (Check)*
 - *Actuar (Act)*



148

Antes de entrar en detalle sobre las herramientas vamos a pensar un poco acerca de la estrategia de despliegue de IPv6.

Podemos pensar en varios ejemplos históricos de despliegues de nuevas tecnologías en que no se tuvieron en cuenta desde el inicio los aspectos de seguridad, como cuando surgieron los routers inalámbricos. Podemos mencionar varios ejemplos de la vida diaria de proyectos de última hora, demostraciones para clientes, que terminan por convertirse en sistemas en producción y presentan diferentes vulnerabilidades.

Lo ideal es que el despliegue de IPv6 se realice de una manera planificada y organizada, tomando en cuenta la seguridad desde el primer momento.

Este slide y los siguientes fueron tomados de la presentación que realizó Joe Klein en la Google IPv6 Implementors Conference 2009. El original se encuentra en: https://sites.google.com/site/ipv6implementors/conference2009/agenda/03_Klein_IPv6_Security.pdf?attredirects=0

Desarrollo de sistemas con IPv6 habilitado

Fecha	Productos	Soporte para IPv6	IPv6 habilitado
1996	OpenBSD / NetBSD / FreeBSD	Sí	Sí
	Linux Kernel 2.1.6	Sí	No
1997	AIX 4.2	Sí	No
2000	Windows 95/98/ME/NT 3.5/NT 4.0	Sí (paquetes adicionales)	No
	Windows 2000	Sí	No
	Solaris 2.8	Sí	Sí
2001	Cisco IOS (12.x y superior)	Sí	No
2002	Juniper (5.1 y superior)	Sí	La mayoría
	IBM z/OS	Sí	Sí
	Apple OS/10.3	Sí	Sí
	Windows XP	Sí	No
	Linux Kernel 2.4	Sí	No
	AIX 6	Sí	Sí
	IBM AS/400	Sí	Sí
2006	Routers Linksys (Mindspring)	Sí	No
	Teléfonos celulares (varios)	Sí	Sí
	Solaris 2.10	Sí	Sí
	Linux Kernel 2.6	Sí	Sí
2007	Apple Almost Extreme	Sí	Sí

Es interesante observar cuántos sistemas son capaces de ejecutar IPv6, y que muchos de ellos vienen con el protocolo habilitado por defecto. La lista incluye sistemas operativos, teléfonos celulares, equipos de red, entre otros.

Incidentes de seguridad en IPv6

2001	Revisión de logs, anuncio del Proyecto Honeynet
2002	Proyecto Honeynet: Lance Spitzner: Solaris Snort: Martin Roesch: IPv6 se agregó, luego se eliminó
2003	Gusano: W32.HLLW.Raleka: Descarga de archivos de un local predefinido y conexión a un IRC server
2005	Troyano: Troj/LegMir-AT: Conexión a un IRC server CERT: Backdoors usando Teredo IPv6 Mike Lynn: Blackhat: Captura de paquetes IPv6
2006	CAMSECWest: THC IPv6 Hacking Tools RP Murphy: DefCon: Backdoors IPv6
2007	Rootkit: W32/Agent.EZM!tr.dldr: TCP HTTP SMTP James Hoagland: Blackhat: Falla informada en el Teredo IPv6 de Vista
2008	HOPE: Vulnerabilidad en teléfonos móviles con IPv6 Noviembre: "Los atacantes lo intentarán utilizar como mecanismo de transporte para botnets. IPv6 se ha vuelto un problema desde el punto de vista operativo." Arbor Networks

150

Se vienen reportando incidentes de seguridad que involucran IPv6 desde hace algún tiempo.

Ejemplo:

*From: Lance Spitzner <lance_at_honeynet.org>
Date: Tue, 17 Dec 2002 20:34:33 -0600 (CST)*

Recently one of the Honeynet Project's Solaris Honeynets was compromised.

What made this attack unique was after breaking into the system, the attackers enabled IPv6 tunneling on the system, with communications being forwarded to another country. The attack and communications were captured using Snort, however the data could not be decoded due to the IPv6 tunneling. Also, once tunneled, this could potentially disable/bypass the capabilities of some IDS systems.

Marty is addressing this issue and has added IPv6 decode support to Snort. Its not part of Snort current (2.0) yet, its still in the process of testing. If you would like to test this new capability, you can find it online at <http://www.snort.org/~roesch/>

Marty's looking for feedback. As IPv6 usage spreads, especially in Asia, you will want to be prepared for it. Keep in mind, even in IPv4 environments (as was our Solaris Honeynet) attackers can encode their data in IPv6 and then tunnel it through IPv4. We will most likely being seeing more of this type of behavior.

Just a friendly heads-up :)

-- Lance Spitzner <http://www.tracking-hackers.com>

Malware

	Fecha	Infección	Nombre
2001	10/1/2001	DOS bot	Ipv4.ipv6.tcp.connection
2003	9/26//2003	Gusano	W32/Raleka!worm
2004	7/6/2004	Gusano	W32/Sdbot-JW
2005	2/18/2005	Gusano	W32/Sdbot-VJ
	8/24/2005	Troyano	Troj/LegMir-AT
	9/5/2005	Troyano	Troj/LegMir-AX
2006	4/28/2006	Troyano	W32/Agent.ABU!tr.dldr
2007	1/2/2007	Troyano	Cimuz.CS
	4/10/2007	Troyano	Cimuz.EL
	5/4/2007	Troyano	Cimuz.FH
	11/5/2007	Gusano	W32/Nofupat
	11/15/2007	Troyano	Trojan.Astry
	12/1/2007	Rootkit	W32/Agent.EZM!tr.dldr
	12/16/2007	Troyano	W32/Agent.GBU!tr.dldr
	12/29/2007	Gusano	W32/VB-DYF
2008	4/22/2008	Troyano	Troj/PWS-ARA

151

Del mismo modo, se viene reportando *malware* que utiliza IPv6 o ataca sistemas IPv6 al menos desde 2001.

Vulnerabilidades IPv6



152

cgi.br

El número de vulnerabilidades (encontradas) crecerá a medida que aumente el uso de IPv6.

Impactos de las vulnerabilidades

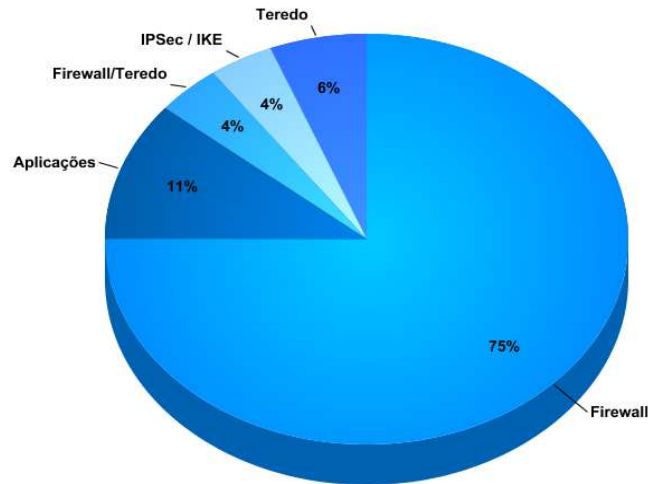
Vulnerabilidades IPv6 publicadas por clasificación



La mayor parte de las vulnerabilidades expone los sistemas a ataques DoS.

Problemas centrais

Vulnerabilidades IPv6 publicadas por tecnologia

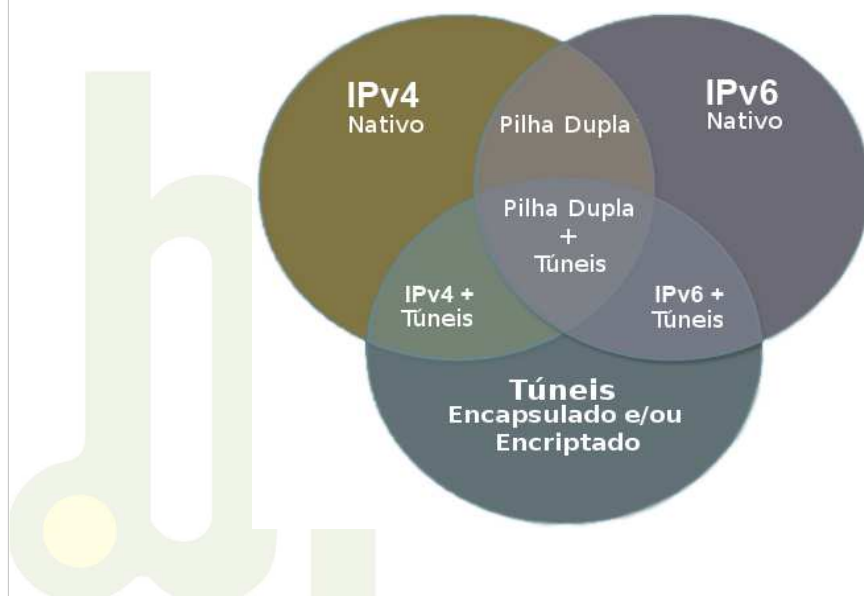


154

egi.br

La mayor parte de las vulnerabilidades publicadas afecta los equipos de red

Áreas de ataque



155

egi.br

Durante la transición de IPv4 a IPv6 se utilizan ambas tecnologías en forma nativa y túneles. Observando la figura se pueden ver 7 superficies de ataque posibles en este contexto.

Blanco de 7 capas



156

Aunque la IP corresponde a la capa de red, sus consecuencias en otras capas también pueden generar vulnerabilidades o problemas.

Seguridad en IPv6



Copyright © 2003 David Farley, d-farley@ibiblio.org
<http://ibiblio.org/Dave/drfun.html>
 This cartoon is made available on the Internet for personal viewing only. Opinions expressed herein are solely those of the author.

The brave new world of IPv6

157

cgi.br

IPv6 ofrece diferentes herramientas tanto para defensa como para ataque:

Defensa:

- IPsec
- SEND
- *Crypto-Generated Address*
- *Unique Local Addresses*
- *Privacy Addresses*

Ataque:

- Túnel automático
- Descubrimiento de vecinos y autoconfiguración
- Modelo end-to-end
- Novedad / Complejidad
- Falta de políticas, capacitación y herramientas.

Es necesario:

- Preocuparse por la seguridad y considerar la seguridad de los equipos desde el inicio
- Obtener equipos certificados
- Educación / Capacitación
- Actualizar las herramientas y procesos de seguridad
- Desarrollar prácticas de programación adecuadas (y seguras) para IPv6
- Buscar auditores / equipos de prueba que conozcan IPv6

IPSec

- Implementa encriptación y autenticación de paquetes en la capa de red.
- Proporciona una solución de seguridad *end-to-end*
 - Asociaciones de seguridad
- Garantiza la integridad, confidencialidad y autenticidad de los datos.
- Desarrollado como parte integral de IPv6.
 - Soporte obligatorio.
- Adaptado para funcionar con IPv4.
 - Soporte opcional.

158

IPSec fue desarrollado para IPv4 y es poco lo que cambia con IPv6. No obstante, el soporte pasa a ser obligatorio y no hay NAT para entorpecer el funcionamiento.

IPSec - Modos de Operação

- IPSec puede operar en dos modos diferentes:

Modo Transporte

Cabeçalho IP original	Cabeçalho IPsec	TCP	Dados
-----------------------	-----------------	-----	-------

|----- Pode ser encriptado -----|

Modo Túnel (VPN de Capa 3)

Novo Cab. IP	Cabeçalho IPsec	Cab. IP original	TCP	Dados
--------------	-----------------	------------------	-----	-------

|..... Pode ser encriptado|

159

IPSec se puede utilizar de dos maneras diferentes: modo transporte o modo túnel. En modo transporte, ambos extremos de la comunicación requieren soporte IPSec para que entre ellos la comunicación sea segura.

A diferencia de lo que ocurre en modo transporte, en modo túnel (también conocido como VPN) IPSec se implementa en dispositivos propios (por ejemplo, concentradores VPN) entre los que la comunicación IPSec se realiza encapsulando todos los paquetes IP de los respectivos extremos.

Cabe destacar que en el caso de una comunicación en modo transporte se mantiene el encabezado del paquete IP original. En modo túnel éste se codifica y se crea un nuevo encabezado que hace posible la comunicación entre el dispositivo emisor y el dispositivo receptor (del túnel).

- **Transporte** - Protege solo los protocolos de las capas superiores, ya que el encabezado de seguridad aparece inmediatamente después del encabezado IP y antes de los encabezados de los protocolos de las capas superiores;
- **Túnel** - Protege todo el paquete IP, encapsulándolo dentro de otro paquete IP y dejando visible solo el encabezado IP externo.

IPSec

- *Framework* de seguridad - Utiliza recursos independientes para realizar sus funciones.
 - *Authentication Header (AH)*
 - Integridad de todo el paquete;
 - Autenticación del origen;
 - Protección contra el reenvío del paquete.
 - *Encapsulating Security Payload (ESP)*
 - Confidencialidad;
 - Integridad del interior del paquete;
 - Autenticación del origen;
 - Protección contra el reenvío del paquete.
 - *Internet Key Exchange (IKE)*
 - Generar y administrar claves de seguridad.

IPSec - AH

- *Authentication Header (AH)*

Siguiente encabezado	Tam. encab. de extensión	Reservado
Índice de parámetros de seguridad		
Número de secuencia		
Autenticación de los datos		

- Se agrega después de los encabezados *Hop-by-Hop*, *Routing* y *Fragmentation* (si corresponde);
- Se puede utilizar en ambos modos de operación.

IPSec - ESP

- *Encapsulating Security Payload (ESP)*

Índice de parámetros de seguridad		
Número de secuencia		
Datos + Relleno		
	Tamaño de relleno	Siguiente encabezado
Autenticación de los datos		

- Responsable por la encriptación de los datos (opcional);
- Se puede utilizar en ambos modos de operación;
- Se puede combinar con AH.

IPSec - Administración de Claves

- Manual
 - Claves configuradas en cada sistema.
- Automática
 - *Internet Key Exchange* (IKE)
 - Se basa en tres protocolos
 - ISAKMP
 - OAKLEY
 - SKEME
 - Funciona en dos fases
 - Tiene dos versiones
 - IKEv1
 - IKEv2

163

La Administración de Claves es una de las dificultades operativas para la implantación de IPSec en IPv4, y continúa presentando el mismo nivel de complejidad que en IPv6.

Más información:

- RFC 4301 - *Security Architecture for the Internet Protocol*

SEcure Neighbor Discovery - SEND

- IPv4 - Ataques al ARP y al DHCP (Spoofing).
 - No hay mecanismos de protección.
- IPv6 - Utiliza el protocolo de Descubrimiento de Vecinos.
 - Mensajes ICMPv6 - No depende de la capa de enlace;
 - Tiene las mismas vulnerabilidades que el ARP y el DHCP;
 - Dificultadas para la implementación de IPSec.
 - Problemas en la generación automática de claves.

Más información:

- RFC 3756 - *IPv6 Neighbor Discovery (ND) Trust Models and Threats*
- RFC 3971 - *SEcure Neighbor Discovery (SEND)*

SEND

- Cadena de certificados.
 - Se utilizan para certificar la autoridad de los routers.
- Utiliza direcciones CGA.
 - Generadas criptográficamente.
- Nueva opción del protocolo de Descubrimiento de Vecinos.
 - *RSA signature* - Protege los mensajes relativos a *Neighbor Discovery* y *Router Discovery*.
- Dos nuevas opciones del protocolo de Descubrimiento de Vecinos.
 - *Timestamp* y *Nonce* - Previene ataques de reenvío de mensajes.

165

Estructura de las direcciones

- Los 128 bits de espacio de direccionamiento pueden dificultar algunos tipos de ataques.
- Nuevas formas de generar IID.
- También cambia el filtrado de direcciones.
 - Direcciones *bogons*.
- Nuevos tipos de ataques.

Estructura de las direcciones

- Barrido de direcciones (*scanning*)
 - Se vuelve más complejo, pero no imposible.
 - Con una máscara estándar /64 son posibles 2^{64} direcciones por subred.
 - Recorriendo 1 millón de direcciones por segundo, se necesitarían más de 500.000 años para recorrer toda la subred.
 - NMAP solo tiene soporte para escanear un único *host* por vez.
 - Los gusanos que utilizan esta técnica para infectar otros dispositivos también tendrán dificultades para continuar propagándose.

167

Estructura de las direcciones

- Barrido de direcciones (*scanning*)
 - Deben surgir nuevas técnicas:
 - Explorar direcciones de servidores públicos anunciadas en el DNS.
 - Búsqueda de direcciones fáciles de memorizar utilizadas por los administradores de redes.
 - ::10, ::20, ::DAD0, ::CAFE.
 - Último byte de la dirección IPv4.
 - Explorar direcciones asignadas automáticamente con base en la MAC, fijando del número correspondiente al fabricante de la placa de red.

168

Direcciones - CGA

- Direcciones IPv6 cuyas IID se generan criptográficamente utilizando una función hash de claves públicas.
 - Prefijo /64 de la subred.
 - Clave pública del propietario de la dirección.
 - Parámetro de seguridad.
- Utiliza certificados X.509.
- Utiliza la función hash SHA-1.

Más información:

- RFC 3972 - *Cryptographically Generated Addresses (CGA)*

Direcciones - Extensiones de privacidad

- Extensión del mecanismo de autoconfiguración *stateless*.
- Genera direcciones temporarias y/o aleatorias.
- Dificulta el rastreo de dispositivos o usuarios.
- Las direcciones cambian de acuerdo con la política local.
- Para cada dirección generada se debe ejecutar la Detección de Direcciones Duplicadas.

170

Más información:

- RFC 4864 - *Local Network Protection for IPv6*
- RFC 4941 - *Privacy Extensions for Stateless Address Autoconfiguration in IPv6*

Seguridad en IPv6

- La seguridad de las redes IPv6 no difiere sustancialmente de la seguridad de las redes IPv4.
- Muchas formas de ataque continúan iguales, así como la manera de evitarlas.
 - *Sniffing*
 - Ataques a la capa de aplicación
 - *Man-in-the-Middle*
 - Virus
 - DoS
- IPSec no es la solución a todos los problemas.

171

cgi.br

Muchas implementaciones de pila IPv6 todavía no soportan IPSec en forma integral. Por lo tanto, éste se está usando sin soporte criptográfico. Pero aunque se esté utilizando, en las redes IP todavía preocupan varios temas de seguridad. No obstante lo anterior, IPv6 potencialmente puede mejorar la seguridad en Internet.

- DoS: En IPv6 no existen direcciones *broadcast*
 - Evita ataques mediante el envío de paquetes ICMP a la dirección *broadcast*.
- Las especificaciones de IPv6 prohíben la generación de paquetes ICMPv6 en respuesta a mensajes enviados a direcciones globales *multicast* (excepto mensajes *packet too big*).
 - Muchos sistemas operativos siguen la especificación;
 - Todavía hay cierta incertidumbre sobre el riesgo que pueden crear los paquetes ICMPv6 que se originan en direcciones *multicast* globales.

Recomendaciones

- Implementar extensiones de privacidad solo en las comunicaciones externas.
- Cuidado con el uso indiscriminado, puede dificultar las auditorías internas.
- Las direcciones de uso interno se deben filtrar en los routers de borde.
- Las direcciones *multicast* como **FF02::1** (*all-nodes on link*), **FF05::2** (*all-routers on link*) y **FF05::5** (*all DHCPv6 servers*) se pueden transformar en nuevos vectores de ataque.
- Filtrar el ingreso de paquetes con direcciones de origen *multicast*.

Recomendaciones

- No usar direcciones obvias;
- Filtrar servicios innecesarios en el firewall.
- Filtrar mensajes ICMPv6 no esenciales.
- Filtrar direcciones *bogon*.
 - En IPv6 este filtrado es diferente al que se realiza en IPv4.
 - En IPv4 se bloquean los rangos no asignados (son pocos).
 - En IPv6 es al revés. Es más fácil liberar solo los rangos asignados.

173

Más información:

- RFC 3704 - *Ingress Filtering for Multihomed Networks*
- RFC 4890 - *Recommendations for Filtering ICMPv6 Messages in Firewalls*

Recomendaciones

- Bloquear fragmentos de paquetes IPv6 con destino a equipos de red.
- Descartar paquetes de tamaño menor a 1280 bytes (excepto el último).
- Los mecanismos de seguridad de BGP y de IS-IS no cambian.
- Con OSPFv3 y RIPng se debe utilizar IPSec.
- Limitar el número de saltos para proteger dispositivos de red.
- Utilizar IPSec siempre que sea necesario.

IPv6.br

La nueva generación del
Protocolo de Internet

175

egi.br

Coexistencia y transición

Módulo 7

176

cgi.br

Para que la transición entre ambos protocolos se produzca de forma gradual y sin mayores impactos sobre el funcionamiento de las redes es necesario que exista un período de coexistencia entre los protocolos IPv4 e IPv6.

En este módulo conoceremos las diferentes técnicas de transición utilizadas, analizando los conceptos básicos del funcionamiento de la doble pila, los túneles y las traducciones, para así comprender en qué situaciones es mejor aplicar cada una de ellas.

Coexistencia y transición

- Toda la estructura de Internet está basada en IPv4.
- Un cambio inmediato de protocolo es inviable debido al tamaño y a la proporción que tiene esta red.
- La adopción de IPv6 se debe realizar de manera gradual.
- Inicialmente habrá un período de transición y de coexistencia entre los dos protocolos.
- Las redes IPv4 necesitarán comunicarse con las redes IPv6 y viceversa.
- Para facilitar este proceso se han desarrollado algunas técnicas que buscan mantener la compatibilidad de toda la base instalada de redes IPv4 con el nuevo protocolo IPv6.

177

Con el objetivo de facilitar el proceso de transición entre las dos versiones del Protocolo de Internet se han desarrollado algunas técnicas para que toda la base de redes instaladas sobre IPv4 se mantenga compatible con IPv6, ya que en este primer momento de coexistencia de los dos protocolos esta compatibilidad es fundamental para el éxito de la transición a IPv6.

Cada una de estas técnicas tiene características específicas y se puede utilizar individualmente o junto con otras técnicas para adecuarse a las necesidades de cada situación, ya sea una migración a IPv6 que se realiza paso a paso, comenzando por un único *host* o subred, o incluso toda una red corporativa.

Coexistencia y transición

- Estas técnicas de transición se dividen en 3 categorías:
 - **Doble pila**
 - Provee soporte a ambos protocolos en el mismo dispositivo.
 - **Tunelización**
 - Permite el tráfico de paquetes IPv6 sobre la estructura de la red IPv4 existente.
 - **Traducción**
 - Permite la comunicación entre nodos que solo soportan IPv6 y nodos que solo soportan IPv4.

178

cgi.br

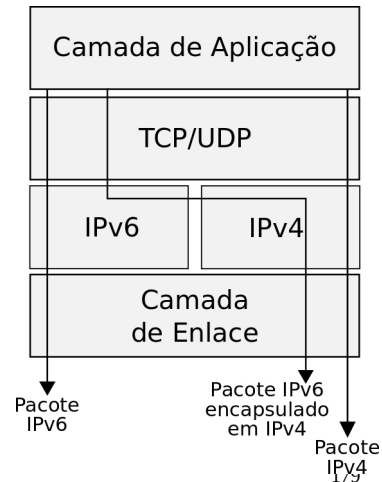
Estos mecanismos de transición se pueden clasificar en las siguientes categorías:

- **Doble pila:** Provee soporte a ambos protocolos en el mismo dispositivo;
- **Tunelización:** Permite el tráfico de paquetes IPv6 sobre estructuras de red IPv4; y
- **Traducción:** Permite la comunicación entre nodos que solo soportan IPv6 y nodos que solo soportan IPv4.

Como el período de coexistencia de los dos protocolos puede durar indefinidamente, la implementación de métodos que permitan la interoperabilidad entre IPv4 e IPv6 garantizará una migración segura hacia el nuevo protocolo, a través de la realización de pruebas que permitan conocer las opciones que ofrecen dichos mecanismos, y además podrían evitar que en el futuro surjan "islas" de comunicación aisladas.

Doble pila

- Los nodos se vuelven capaces de enviar y recibir paquetes tanto para IPv4 como para IPv6.
- Al comunicarse con un nodo IPv6, un nodo IPv6/IPv4 se comporta como un nodo IPv6; al comunicarse con un nodo IPv4, como un nodo IPv4.
- Necesita al menos una dirección para cada pila.
- Utiliza mecanismos IPv4, como por ejemplo DHCP, para adquirir direcciones IPv4, y mecanismos de IPv6 para direcciones IPv6.



En esta fase inicial de implementación de IPv6 todavía no es aconsejables tener nodos que solamente soporten esta versión del protocolo IP, ya que muchos servicios y dispositivos de red continúan trabajando solamente sobre IPv4. Por este motivo, una posibilidad consiste en implementar el método conocido como doble pila.

El uso de este método permite que los *hosts* y routers estén equipados con pilas para ambos protocolos y tengan la capacidad de enviar y recibir ambos tipos de paquetes, IPv4 e IPv6. Así, en la comunicación con un nodo IPv6 un nodo doble pila (o nodo IPv6/IPv4) se comportará como un nodo solo IPv6, mientras que en la comunicación con un nodo IPv4 se comportará como un nodo solo IPv4.

Cada nodo IPv6/IPv4 se configura con ambas direcciones, utilizando mecanismos IPv4 (por ejemplo DHCP) para adquirir su dirección IPv4 y mecanismos del protocolo IPv6 (por ejemplo autoconfiguración y/o DHCPv6) para adquirir su dirección IPv6.

Este método de transición puede facilitar la gestión de la implementación de IPv6, ya que permite implementar IPv6 en forma gradual, configurando pequeñas secciones del entorno de red cada vez. Además, si en el futuro se dejara de utilizar IPv4, bastaría con deshabilitar la pila IPv4 de cada nodo.

Doble pila

- Una red doble pila es una infraestructura capaz de encaminar ambos tipos de paquetes.
- Exige analizar algunos aspectos:
 - Configuración de los servidores de DNS;
 - Configuración de los protocolos de enrutamiento;
 - Configuración de los *firewalls*;
 - Cambios en la administración de las redes.

180

cgi.br

Al implementar la técnica de doble pila es necesario considerar algunos aspectos. Se debe analizar la necesidad de realizar cambios en la infraestructura de red, como por ejemplo la estructuración del servicio de DNS y la configuración de los protocolos de enrutamiento y los *firewalls*.

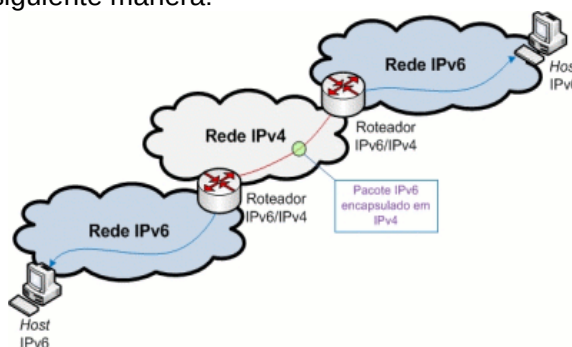
Con respecto al DNS, es necesario que éste esté habilitado para resolver nombres y direcciones de ambos protocolos. En el caso de IPv6, es necesario responder a consultas de registros tipo AAAA (quad-A), que almacenan direcciones en formato IPv6, y al dominio creado para la resolución reversa, ip6.arpa. Si desea obtener más información acerca del soporte IPv6 del DNS, consulte la RFC 3596.

En una red IPv6/IPv4 la configuración del enrutamiento IPv6 normalmente es independiente de la configuración del enrutamiento IPv4. Esto implica que si antes de la implementación de la pila doble la red solo utilizaba el protocolo de enrutamiento interno OSPFv2 (que solo soporta IPv4), será necesario migrar a un protocolo de enrutamiento que soporte tanto IPv6 como IPv4, como por ejemplo IS-IS, o forzar la ejecución de un IS-IS o OSPFv3 en paralelo con el OSPFv2.

La forma en que se realiza el filtrado de paquetes que atraviesan la red puede depender de la plataforma que se esté usando. Por ejemplo, en un entorno Linux los filtros de paquetes son totalmente independientes entre sí, de modo que iptables filtra solamente paquetes IPv4 e ip6tables solo IPv6, sin compartir ninguna configuración. En FreeBSD las reglas se aplican a ambos protocolos, a menos que explícitamente se restrinja a cuál familia de protocolo deben ser aplicadas usando inet o inet6.

Técnicas de tunelización

- También llamado encapsulamiento.
- El contenido del paquete IPv6 se encapsula en un paquete IPv4.
- Se pueden clasificar de la siguiente manera:
 - *Router-a-Router*
 - *Host-a-Router*
 - *Router-a-Host*
 - *Host-a-Host*



La técnica de creación de túneles – tunelización – permite transmitir paquetes IPv6 a través de la infraestructura IPv4 existente sin necesidad de realizar ningún cambio en los mecanismos de enrutamiento, encapsulando el contenido del paquete IPv6 en un paquete IPv4.

Estas técnicas, descritas en la RFC 4213, han sido las más utilizadas en la fase inicial del despliegue de IPv6 por ser fáciles de aplicar en entornos de prueba en los cuales las redes no están estructuradas para ofrecer tráfico IPv6 nativo.

Más información:

- RFC 4213 - *Basic Transition Mechanisms for IPv6 Hosts and Routers*

Técnicas de tunelización

- Existen diferentes formas de encapsulamiento:
 - Paquetes IPv6 encapsulados en paquetes IPv4;
 - Protocolo 41.
 - 6to4, ISATAP y *Tunnel Brokers*.
 - Paquetes IPv6 encapsulados en paquetes GRE;
 - Protocolo GRE.
 - Paquetes IPv6 encapsulados en paquetes UDP;
 - TEREDO.

182

Hay diferentes técnicas de tunelización disponibles: Los escenarios en los que se pueden aplicar, las dificultades de implementación y la diferencia de rendimiento varían significativamente entre los diferentes modelos, por lo que es necesario realizar un análisis detallado de cada uno de ellos. Las principales técnicas de tunelización son las siguientes:

- Tunnel Broker
- 6to4
- ISATAP
- Teredo
- GRE

Ahora analizaremos detalladamente cada una de estas técnicas.

Tunnel Broker

- Consiste en un túnel IPv6 dentro de la red IPv4, que se crea entre su computadora o red y el proveedor que suministrará la conectividad IPv6.
- Basta con registrarse en un proveedor de acceso *Tunnel Broker* y descargar un *software* o *script* de configuración.
- La conexión del túnel se realiza solicitando el servicio al Servidor Web del proveedor.
- Indicado para redes pequeñas o para un único *host* aislado.

183

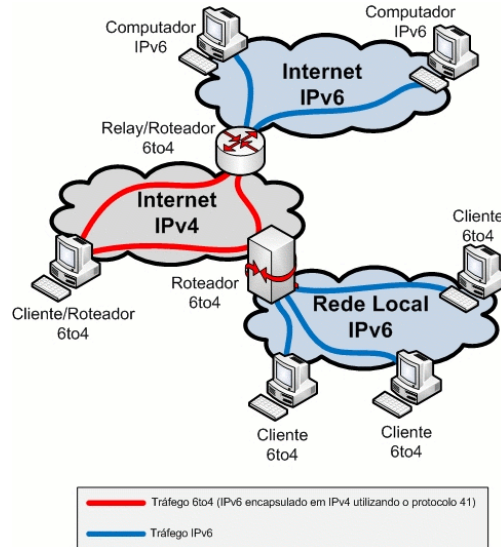
Descrita en la RFC 3053, esta técnica permite que *hosts* IPv6/IPv4 aislados en una red IPv4 accedan a redes IPv6. Su funcionamiento es bastante sencillo. Primero es necesario registrarse en un proveedor de acceso Tunnel Broker y descargar un *software* o *script* de configuración. La conexión del túnel se establece solicitando el servicio al Servidor Web del proveedor, quien luego de una autenticación verifica qué tipo de conexión está utilizando el cliente (IPv4 pública o NAT) y le asigna una dirección IPv6. A partir de ese momento el cliente puede acceder a cualquier *host* en Internet.

Más información:

- RFC 3053 - *IPv6 Tunnel Broker*

6to4

- Forma de tunelización router-a-router.
- Provee una dirección IPv6 única al *host*.
- La dirección está formada por el prefijo de dirección global **2002:wwwx:yyzz::/48**, donde **wwwx:yyzz** es la dirección IPv4 pública del *host* convertida a formato hexadecimal.
- El *relay* 6to4 se identifica mediante la dirección *anycast* **192.88.99.1**.
- Encaminamiento asimétrico.
- Se puede utilizar con relays públicos cuando no hay conectividad v6 nativa.
- Cuando hay conectividad y servicios nativos se debe implementar para facilitar la comunicación con clientes 6to4.



cgi.br

Definida en la RFC 3056, la técnica de tunelización automática 6to4 permite la interconexión punto-a-punto entre routers, subredes o computadoras IPv6 a través de la red IPv4, proporcionando una dirección IPv6 única que se forma a partir de direcciones IPv4 públicas. Este direccionamiento 6to4 utiliza el prefijo de dirección global **2002:wwwx:yyzz::/48**, donde **wwwx:yyzz** es la dirección IPv4 pública del cliente convertida a formato hexadecimal.

* **Cliente/Router 6to4:** Cliente que tiene una dirección IPv4 pública y conectividad directa 6to4, es decir, que tiene una interfaz virtual 6to4 por la cual accede directamente a la Internet IPv6 sin necesidad de utilizar un router 6to4. Solo requiere un *Relay* 6to4;

* **Router 6to4:** Router que soporta 6to4, permitiendo que los clientes que no soportan este tipo de direcciones accedan a otros *hosts* 6to4 IPv6 a través del mismo. En el caso de los accesos a la Internet IPv6, éste direcciona el tráfico hasta el *Relay Router* más cercano, que encaminará el paquete hacia la red IPv6;

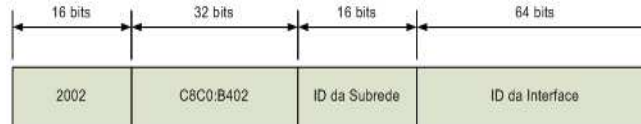
* **Relay 6to4:** Router con soporte para 6to4 y que también tiene conexión nativa a la Internet IPv6. De este modo puede enrutar y comunicarse con la red IPv6 nativa, con la red IPv4 y con la red 6to4;

* **Cliente 6to4:** Equipo de red o computadora que solo tiene dirección IPv6 en formato 6to4, pero que no tiene una interfaz virtual 6to4. Por lo tanto, necesita un router 6to4 para realizar la comunicación con otras redes IPv6 y 6to4.

Más información:

- RFC 3056 - *Connection of IPv6 Domains via IPv4 Clouds*

6to4



- El prefijo 6to4 siempre es **2002**.
- El siguiente campo, IPv4 pública del cliente, se crea convirtiendo la dirección a formato hexadecimal.
- El ID de la subred se puede usar para segmentar la red IPv6 6to4 en hasta 2^{16} subredes con 2^{64} direcciones cada una; se puede utilizar, por ejemplo, 0, 1, 2, 3, 4...
- El ID de la interfaz puede ser igual al segundo campo (Windows lo hace de este modo) o a cualquier otro número en el caso de configuración manual (Linux utiliza numeración secuencial: 1, 2, 3, 4...).

185

- El prefijo 6to4 siempre es **2002**, de acuerdo con la definición de la IANA;
- El siguiente campo, IPv4 pública del cliente, se crea de acuerdo con el siguiente ejemplo:

Dirección IPv4: **200.192.180.002**

Primero convertimos cada número decimal a formato hexadecimal:

200=C8

192=C0

180=B4

002=02

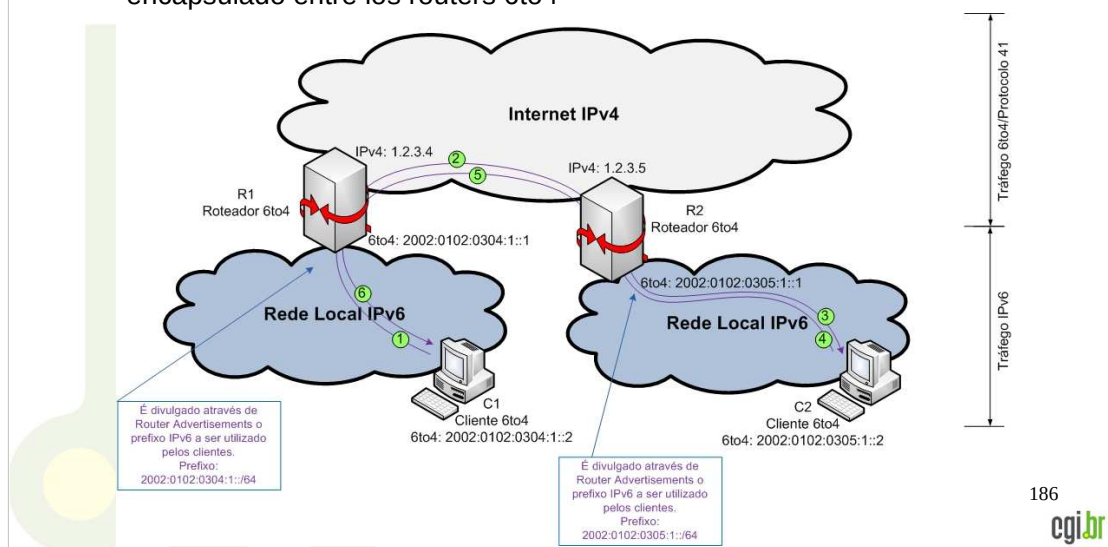
Así la dirección se convierte en C8C0:B402

- El ID de la subred se puede usar para segmentar la red 6to4 asociada a la IPv4 pública en varias subredes(2^{16} subredes con 2^{64} direcciones cada una); puede utilizar, por ejemplo, 0, 1, 2, 3, 4...;
- El ID de la interfaz puede ser igual al segundo campo (IPv4 convertido a formato hexadecimal) en el caso de la configuración automática de Windows Vista y Server 2008, o bien 1, 2, 3, 4... en el caso de configuración manual o de Linux y BSD. Como la longitud de este campo es de 64 bits podemos tener hasta 2^{64} direcciones por cada subred.

6to4

Comunicación entre Clientes 6to4 que están en redes diferentes

- Observe que el tráfico en la red local es nativo IPv6, solo está encapsulado entre los routers 6to4



186

cgi.br

Equipos

Ruta

C1

:::0 a través de R1

2002:0102:0304:1::/64 a través de la interfaz LAN

R1

:::0 a través del Relay 6to4 utilizando la interfaz virtual 6to4

2002::/16 a través de la interfaz virtual 6to4

2002:0102:0304:1/64 hacia la red local a través de la interfaz LAN

R2

:::0 a través de R2

2002:0102:0305:1/64 hacia la red local a través de la interfaz LAN

C2

:::0 a través del Relay 6to4 utilizando la interfaz virtual 6to4

2002::/16 a través de la interfaz virtual 6to4

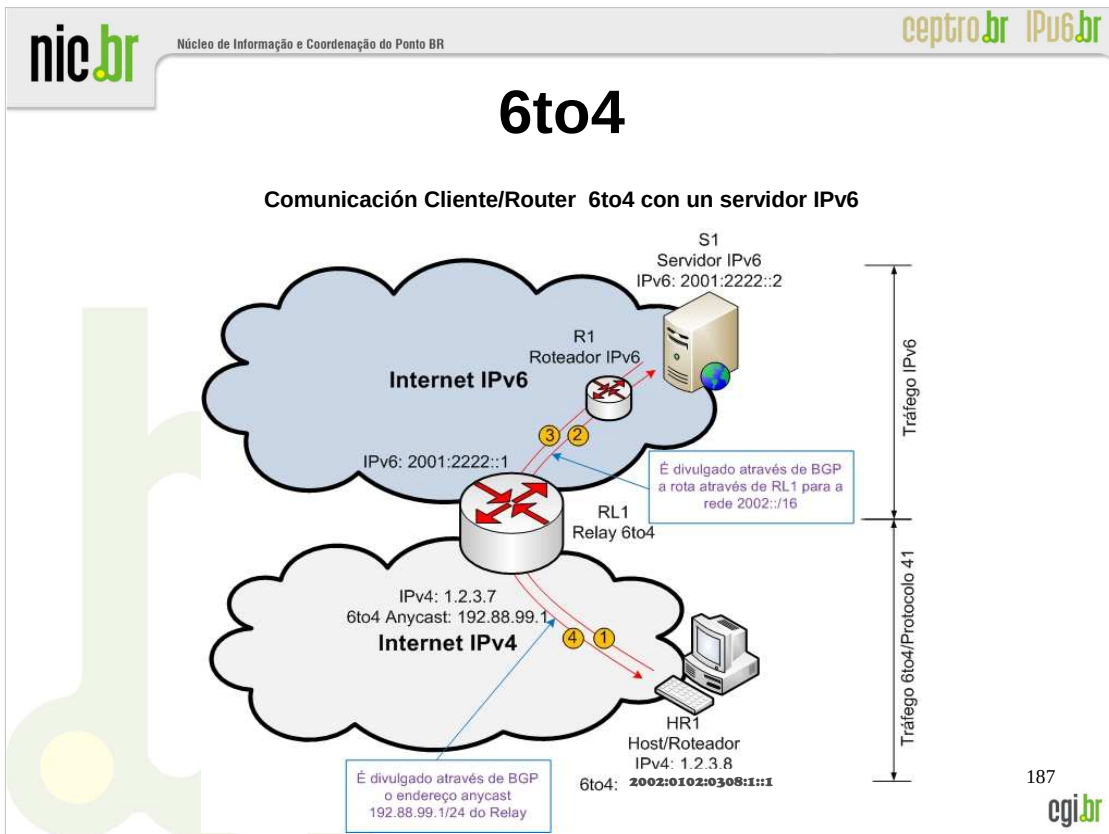
2002:0102:0305:1/64 hacia la red local a través de la interfaz LAN

1- Analizando la tabla de enrutamiento observamos que el paquete se envía a través de la red local IPv6 hacia el router R1 utilizando la ruta :::0;

2- El paquete IPv6 es recibido por R1 a través de la interfaz LAN. R1 verifica su tabla de enrutamiento y descubre que debe enviar el paquete hacia su interfaz virtual 6to4 (ruta a la red **2002::/16**). En esta interfaz el paquete IPv6 se encapsula en un paquete IPv4 (protocolo tipo 41) que se direcciona hacia el router R2 (dirección extraída de la dirección IPv6 del destinatario del paquete original);

3- El paquete IPv6 encapsulado en IPv4 es recibido por R2 a través de su interfaz IPv4 o WAN. Como el paquete es de tipo 41, éste es encaminado a la interfaz 6to4, que lo desencapsula. Al consultar su tabla de enrutamiento R2 descubre que el paquete está destinado a su red local **2002:0102:03:05:1::/64**, por lo que encamina el paquete IPv6 a la computadora C2 a través de su red local.

En los pasos siguientes el sistema de comunicación es el mismo, ya que lo único que cambia es la dirección de encaminamiento del paquete.



Equipo Ruta

HR1 :::0 a través de la interfaz virtual 6to4
 2002::/16 a través de la interfaz virtual 6to4

RL1 :::0 red IPv6 a través de la interfaz LAN
 2002::/16 a través de la interfaz virtual 6to4

S1 Ruta por defecto a través de R1

R1 **2002::/16** a través del Relay RL1 (ruta descubierta a través del anunci vía BGP)p

1- HR1 envía un paquete IPv6 a S1, a través de la tabla de enrutamiento el paquete es direccionado a la interfaz virtual 6to4. Ésta encapsula el paquete IPv6 en un paquete IPv4 (protocolo 41) y coloca como destino la dirección del Relay, que se puede especificar manualmente o descubrir automáticamente encaminando el paquete a la dirección *anycast* **192.88.99.1**;

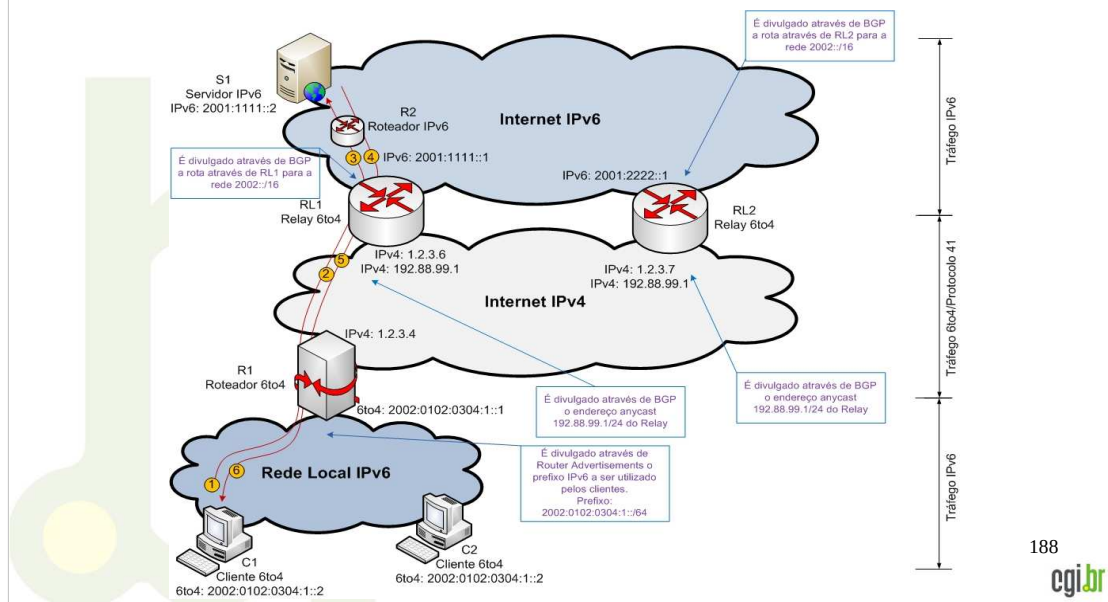
2- El Relay RL1 recibe el paquete encapsulado a través de su IPv4 o *anycast*; como el protocolo del paquete es 41, éste desencapsula el paquete IPv6 y, a través de su tabla de enrutamiento, descubre que el paquete debe ser enviado a S1 a través de su interfaz LAN en la red IPv6;

3- Una vez recibido el paquete, S1 responde utilizando su ruta por defecto a través del router R1 de su red. El router R1 recibe vía BGP la ruta hacia la red **2002::/16** y encamina el paquete al Relay RL1;

4- RL1 recibe el paquete y observa que está destinado a la red 6to4, por lo que encamina el paquete a la interfaz virtual 6to4, que lo encapsula en un paquete IPv4 (protocolo 41) y, a través de la dirección IPv4 implícita en la dirección IPv6 del destinatario, el paquete es encaminado a HR1. HR1 recibe el paquete en su interfaz IPv4, ve que se está utilizando el o protocolo 41 y desencapsula el paquete IPv6 a través de la interfaz virtual 6to4.

6to4

Comunicación Cliente 6to4 con un servidor IPv6 utilizando solo un Relay 6to4 (rutas de ida y vuelta iguales):



Equipo Ruta

RL1	:::0 red IPv6 a través de la interfaz LAN / 2002::16 a través de la interfaz virtual 6to4
RL2	:::0 red IPv6 a través de la interfaz LAN / 2002::16 a través de la interfaz virtual 6to4
S1	Ruta por defecto a través de R2
R2	2002::16 a través del Relay RL1 (ruta descubierta a través del anuncio vía BGP)
R1	:::0 a través del Relay 6to4 RL1 o RL2 utilizando la interfaz virtual 6to4 2002::16 a través de la interfaz virtual 6to4 2002:0102:0304:1/64 hacia la red local a través de la interfaz LAN
C1	:::0 a través de R1 / 2002:0102:0304:1::/64 a través de la interfaz LAN
C2	:::0 a través de R1 / 2002:0102:0304:1::/64 a través de la interfaz LAN

1- De acuerdo con la tabla de enrutamiento, el paquete se envía a través de la red local IPv6 hacia el router R1 utilizando la ruta :::0;

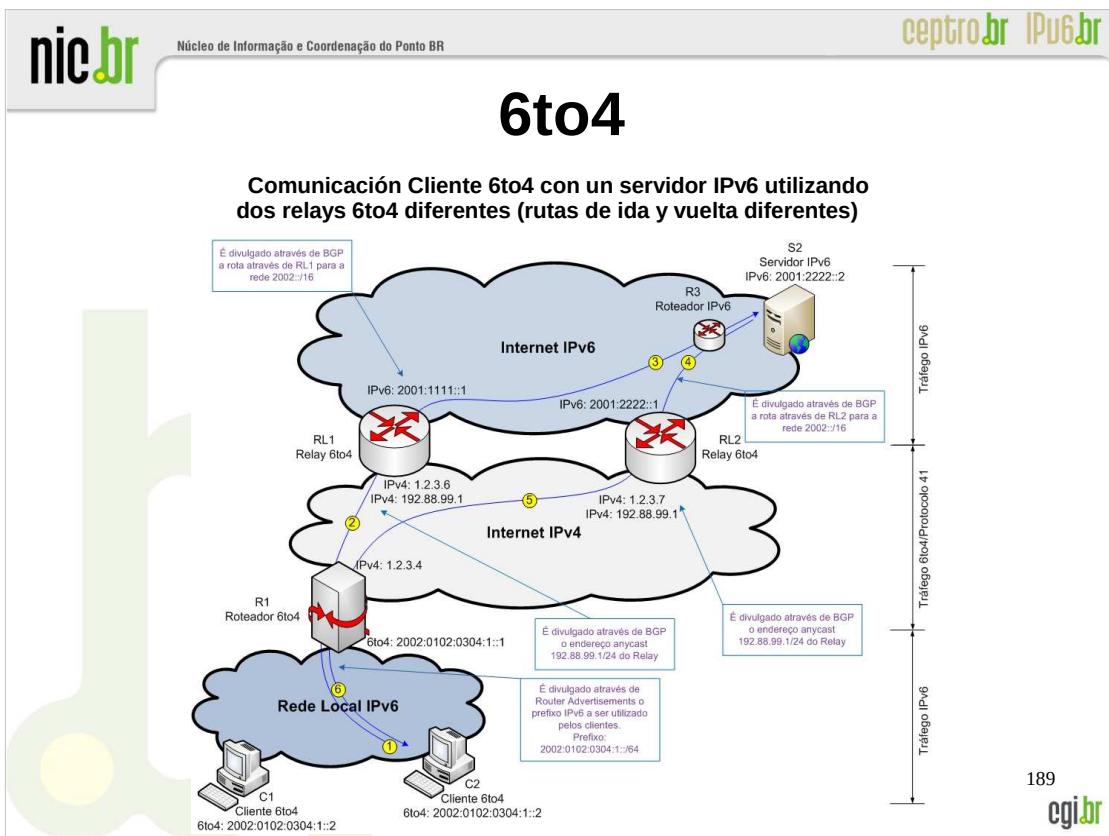
2- El paquete IPv6 es recibido por R1 a través de la interfaz LAN, que verifica su tabla de enrutamiento y descubre que el paquete debe ser encaminado a la interfaz virtual 6to4 (ruta para la red 2002::16). En esta interfaz el paquete IPv6 se encapsula en un paquete IPv4 (protocolo tipo 41) y se envía al Relay RL1 o RL2 (el Relay 6to4 puede ser definido manualmente en el router 6to4 o automáticamente utilizando la dirección *anycast* 192.88.99.1). Supongamos que el paquete se envía al Relay RL1;

3- RL1 recibe el paquete 6to4 a través de su interfaz IPv4 y, como el paquete utiliza el protocolo 41, lo encamina a la interfaz virtual, que desencapsula el paquete IPv6 y verifica en la tabla de enrutamiento que debe enviarlo por la interfaz LAN a través del router R2, que simplemente reenvía el paquete IPv6 al servidor S1;

4- S1 responde enviando otro paquete IPv6 con destino al Cliente C1 utilizando su ruta por defecto que apunta hacia el router R2. R2 recibe el paquete a través de la ruta recibida vía BGP, y sabe que debe enviarlo al relay más próximo, que en este caso es RL1;

5- RL1 recibe el paquete IPv6 y verifica que el destino es la red 6to4 (2002::16). Siendo así, de acuerdo con su tabla de enrutamiento el paquete es encaminado a la interfaz virtual 6to4, que lo empaqueta en un paquete IPv4 (protocolo 41) y lo envía a la dirección IPv4 implícita en la dirección IPv6 del destinatario del paquete;

6- El router R1 recibe el paquete a través de su dirección IPv4 y, como el paquete está utilizando el protocolo 41, éste es encaminado a la interfaz virtual 6to4, que lo desencapsula y verifica la dirección de destino. De acuerdo con su tabla de enrutamiento ésta envía el paquete IPv6 a través de su interfaz LAN al Cliente 6to4 C1.



Equipo

Ruta

RL1	:::0 red IPv6 a través de la interfaz LAN / 2002::/16 a través de la interfaz virtual 6to4
RL2	:::0 red IPv6 a través de la interfaz LAN / 2002::/16 a través de la interfaz virtual 6to4
S2	Ruta por defecto a través de R3
R3	2002::/16 a través del Relay 6to4 RL2 (ruta descubierta a través del anuncio vía BGP)
R1	:::0 a través del Relay 6to4 RL1 o RL2 utilizando la interfaz virtual 6to4 2002::/16 a través de la interfaz virtual 6to4 2002:0102:0304:1/64 hacia la red local a través de la interfaz LAN
C1	:::0 a través de R1 / 2002:0102:0304:1::/64 a través de la interfaz LAN
C2	:::0 a través de R1 / 2002:0102:0304:1::/64 a través de la interfaz LAN

1- De acuerdo con la tabla de enrutamiento, el paquete se envía a través de la red local IPv6 hacia el router R1 utilizando la ruta :::0;

2- El paquete IPv6 es recibido por R1 a través de la interfaz LAN, que verifica su tabla de enrutamiento y descubre que el paquete debe ser enviado a la interfaz virtual 6to4 (ruta para la red 2002::/16). En esta interfaz el paquete IPv6 se encapsula en un paquete IPv4 (protocolo tipo 41) y se envía al Relay RL1 o RL2 (el Relay 6to4 puede ser definido manualmente en el router 6to4 o automáticamente utilizando la dirección *anycast* 192.88.99.1). Supongamos que el paquete se envía al Relay RL1;

3- RL1 recibe el paquete 6to4 a través de su interfaz IPv4, ve que el paquete utiliza el protocolo 41 y lo encamina a la interfaz virtual. Ésta desencapsula el paquete IPv6 y verifica en su tabla de enrutamiento que debe enviarlo por la interfaz LAN a través del router R3, que simplemente reenvía el paquete IPv6 al servidor S2;

4- S2 responde enviando otro paquete IPv6 con destino al Cliente C2 utilizando su ruta por defecto que apunta hacia el router R3. R3 recibe el paquete a través de la ruta recibida vía BGP, y sabe que debe enviarlo al relay más próximo, que es RL2;

5- RL2 recibe el paquete IPv6 y verifica que el destino es la red 6to4 (2002::/16). Siendo así, de acuerdo con su tabla de enrutamiento encamina el paquete a la interfaz virtual 6to4, que lo empaqueta en un paquete IPv4 (protocolo 41) y lo envía a la dirección IPv4 implícita en la dirección IPv6 del destinatario del paquete;

6- El router R1 recibe el paquete a través de su dirección IPv4, verifica que el paquete está utilizando el protocolo 41 y lo encamina a la interfaz virtual 6to4. Ésta lo desencapsula y verifica la dirección de destino. De acuerdo con su tabla de enrutamiento y la dirección de destino, el paquete IPv6 es enviado a través de su interfaz LAN al Cliente 6to4 C2.

6to4

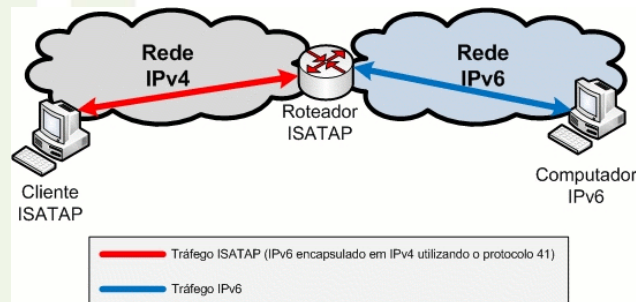
- **Seguridad**

- Los routers relay no verifican los paquetes IPv6 que están encapsulados en IPv4, a pesar de que sí los encapsula y desencapsulan;
- El spoofing de direcciones es un grave problema de los túneles 6to4 que puede ser fácilmente explorado;
- No hay un sistema de autenticación entre el router y el router relay, lo que facilita la exploración de la seguridad mediante el uso de routers relay falsos.

190

ISATAP

- ISATAP (*Intra-Site Automatic Tunnel Addressing Protocol*) - Técnica de tunelización que une *hosts*-a-routers.
- No hay un servicio público de ISATAP, es una técnica que se utiliza dentro de las organizaciones.
- Tiene sentido, por ejemplo, cuando la organización ya tiene numeración IPv6 válida y conectada en el borde pero su infraestructura interna no soporta IPv6.



191

egi.br

La técnica de transición *Intra-Site Automatic Tunnel Addressing Protocol* (ISATAP), definida en la RFC 5214, se basa en túneles IPv6 creados automáticamente dentro de la red IPv4 y en direcciones IPv6 asociadas a esos clientes de acuerdo con el prefijo especificado en el router ISATAP y la dirección IPv4 del cliente. Para crear estos túneles se utilizan las especificaciones de la sección 3 de la RFC 4213, que trata la tunelización a través del protocolo IPv4 tipo 41 o 6in4.

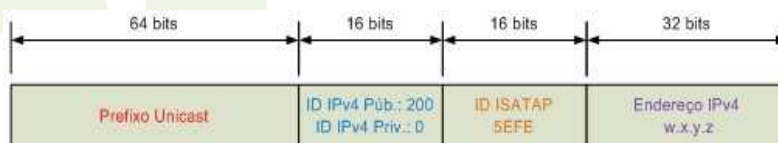
Más información:

- RFC 5214 - *Intra-Site Automatic Tunnel Addressing Protocol* (ISATAP)

ISATAP

• Direccionamiento

- En esta técnica las direcciones IPv4 de los clientes y routers se utilizan como parte de las direcciones ISATAP. Así, un nodo ISATAP puede determinar fácilmente los puntos de entrada y salida de los túneles IPv6 sin utilizar ningún protocolo o recurso auxiliar.
- Las direcciones ISATAP tienen el siguiente formato:



- **Prefijo unicast** : Cualquier prefijo unicast válido en IPv6, puede ser link-local (FE80::/64) o global;
- **ID IPv4 pública o privada**: Si la dirección IPv4 es pública este campo debe tener el valor "200", si es privada (192.168.0.0/16, 172.16.0.0/12 y 10.0.0.0/8) el valor del campo es cero;
- **ID ISATAP**: Siempre tiene el valor 5EFE;
- **Dirección IPv4**: IPv4 del cliente o router en formato IPv4;

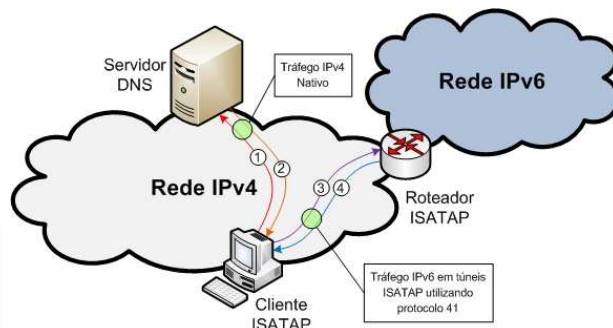
192

Ejemplos de direcciones ISATAP:

Dirección IPv4	Dirección IPv6/ISATAP
250.140.80.1	2001:10fe:0:8003:200:5efe:250.140.80.1 fe80::200:5efe:250.140.80.1
192.168.50.1	2001:10fe:0:8003:0:5efe:192.168.50.1 fe80:0:5efe:250.140.80.1

ISATAP

Início



- 1- Consulta al DNS (en el caso de Windows, busca por ISATAP.dominio-local)
- 2- El servidor de DNS regresa la IPv4 del router ISATAP
- 3- *Router Solicitation* (encapsulado en v4)
- 4- *Router Advertisement* (encapsulado en IPv4)

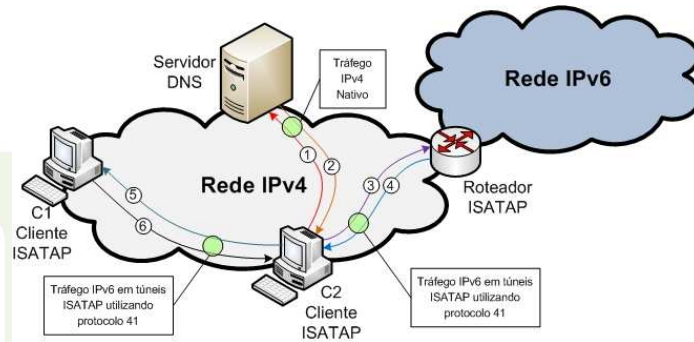
193

egi.br

- 1- En este paso el cliente intenta determinar la dirección IPv4 del router, si la dirección IPv4 todavía no está determinada en su configuración. En el caso de Windows, por defecto intenta resolver el nombre ISATAP e ISATAP.dominio-local mediante resolución local o servidor DNS;
- 2- El servidor de DNS regresa la IPv4 del router ISATAP (si corresponde);
- 3- El cliente envía un mensaje Router Solicitation (RS) encapsulado en IPv4 al router ISATAP;
- 4- El router ISATAP responde con un mensaje Router Advertisement (RA) encapsulado en IPv4, con eso el cliente ya puede configurar sus direcciones IPv6/ISATAP.

ISATAP

Comunicação entre clientes ISATAP en la misma red



- La comunicación entre los clientes ISATAP en una misma red se realiza directamente, sin intervención del router ISATAP (luego de la autoconfiguración inicial). El tráfico en la red es siempre IPv4, IPv6 es encapsulado o desencapsulado localmente en los clientes.

194

egi.br

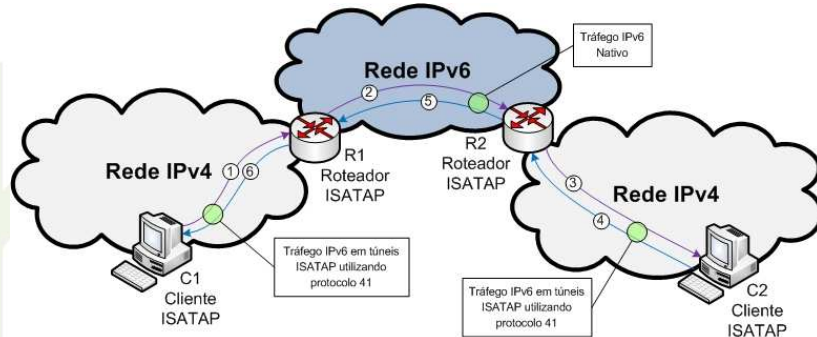
- 1- C2 solicita la resolución DNS del nombre del router ISATAP (si fuera necesario);
- 2- C2 recibe IPv4 del router ISATAP (si fuera necesario);
- 3- El cliente envía un mensaje Router Solicitation (RS) encapsulado en IPv4 al router ISATAP;
- 4- El router ISATAP responde con un mensaje Router Advertisement (RA) encapsulado en IPv4, con eso el cliente ya puede configurar sus direcciones Ipv6/ISATAP;

Los procesos 1 a 4 también son ejecutados por C1;

- 5- El cliente ISATAP C2 envía un paquete IPv6 encapsulado en IPv4 utilizando el protocolo 41 a través de la red IPv4 con destino a la dirección IPv4 de C1;
- 6- El cliente ISATAP C1 recibe el paquete IPv4 y desencapsula el paquete IPv6, luego de lo cual éste responde con otro paquete IPv6 encapsulado en IPv4 utilizando el protocolo 41 a través de la red IPv4 con destino a la dirección IPv4 del cliente C2.

ISATAP

Comunicación entre clientes ISATAP en redes diferentes



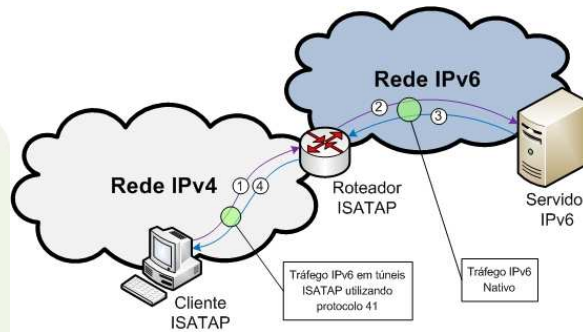
- El tráfico ISATAP entre clientes de redes IPv4 diferentes depende de los routers ISATAP.
- En la red IPv4 el tráfico v6 siempre está encapsulado dentro de paquetes v4.
- Entre routers ISATAP diferentes el tráfico es v6 nativo.

195

1. El cliente ISATAP C1 desea enviar un paquete IPv6 al cliente C2. A través de su tabla de enrutamiento descubre que debe enviarlo utilizando la interfaz virtual ISATAP, por lo que el paquete se encapsula en IPv4 (protocolo 41) y se envía a la dirección IPv4 del router R1;
 2. El router R1 recibe el paquete a través de su interfaz IPv4 pero, como el paquete IPv6 está encapsulado utilizando el protocolo 41, éste lo desencapsula (utilizando la interfaz virtual ISATAP) y verifica la dirección IPv6 del destino. Después de esto descubre que la ruta hacia el destino es a través de la red IPv6, por lo que el paquete desencapsulado (IPv6 nativo) se encamina al router R2;
 3. El router R2 recibe el paquete IPv6 en su interfaz IPv6 pero al verificar la dirección de destino descubre que es para el cliente C2 que está en su subred ISATAP, por lo que envía el paquete a través de esta interfaz, que encapsula nuevamente el paquete IPv6 en un paquete IPv4 y lo envía a C2 (en base a la dirección IPv4 implícita en IPv6). El cliente ISATAP C1 recibe el paquete IPv4 y desencapsula el paquete IPv6 (a través de la interfaz virtual ISATAP);
 4. El cliente ISATAP C2 desea responder al cliente C1, por lo que verifica su tabla de rutas y concluye que debe enviar el paquete a través de la interfaz virtual ISATAP, por lo que el paquete IPv6 es encapsulado en IPv4 y encaminado al router R2;
 5. El router R2 recibe el paquete a través de su interfaz IPv4 pero, como el paquete está utilizando el protocolo 41, éste desencapsula el paquete IPv6 y luego de verificar en su tabla de enrutamiento lo encamina a través de su interfaz IPv6;
 6. El router R1 recibe el paquete IPv6 pero verificando en su tabla de enrutamiento descubre que el paquete debe ser enviado a través de su interfaz virtual ISATAP, la cual encapsula el paquete IPv6 en IPv4 utilizando el protocolo 41 y lo encamina a la dirección IPv4 de C1;
- C1 recibe el paquete pero, como el paquete fue encapsulado utilizando el protocolo 41, éste extrae el paquete IPv6 enviado por C2 y lo recibe.

ISATAP

Comunicación entre clientes ISATAP y computadoras IPv6



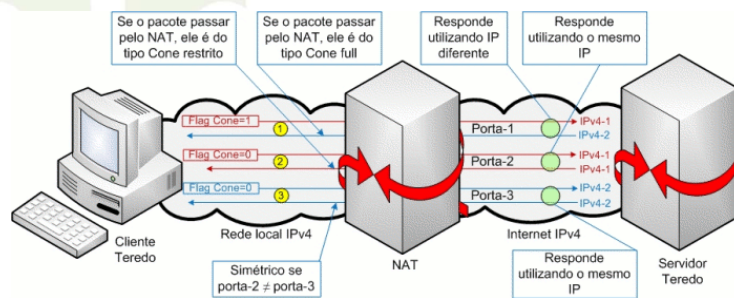
196

egi.br

1. El cliente ISATAP desea enviar un paquete IPv6 al servidor IPv6. A través de su tabla de enrutamiento descubre que debe enviarlo utilizando la interfaz virtual ISATAP, por lo que el paquete se encapsula en IPv4 (protocolo 41) y se envía a la dirección IPv4 del router ISATAP;
2. El router ISATAP recibe el paquete a través de su interfaz IPv4 pero, como el paquete IPv6 está encapsulado utilizando el protocolo 41, éste lo desencapsula (utilizando la interfaz virtual ISATAP) y verifica la dirección IPv6 del destino. Después de esto descubre que la ruta hacia el destino es a través de la red IPv6, por lo que el paquete desencapsulado (IPv6 nativo) se encamina al router IPv6; El servidor recibe el paquete IPv6 destinado al mismo;
- 3- El servidor IPv6 desea responder al cliente ISATAP, por lo que verificando su tabla de enrutamiento descubre que debe enviarlo a través de su ruta por defecto, que es a través de la red IPv6;
- 4- Como la ruta para la red del cliente ISATAP es a través del router ISATAP, el paquete es encaminado al mismo a través de su interfaz IPv6. Verificando en su tabla de enrutamiento el router descubre que debe enviar el paquete a través de su interfaz virtual ISATAP, por lo que el paquete es encapsulado en IPv4 y encaminado al cliente ISATAP a través de la red IPv4. El cliente recibe el paquete IPv4 pero, como está utilizando protocolo 41, desencapsula y recibe el paquete IPv6.

Teredo

- Encapsula el paquete IPv6 en paquetes UDP.
- Funciona a través de NAT tipo cono y cono restringido.
- Envía paquetes *bubbles* periódicamente al servidor para mantener las configuraciones iniciales de la conexión UDP.
- Su funcionamiento es complejo y tiene *overhead*.



197

cgi.br

La técnica de tunelización automática Teredo, definida en la RFC 4380, permite que los nodos ubicados detrás de *Network Address Translations* (NAT) obtengan conectividad IPv6 utilizando el protocolo UDP.

La conexión se realiza a través de un Servidor Teredo que la inicializa y determine el tipo de NAT usado por el cliente. Luego, si el *host* de destino tiene IPv6 nativo, se utiliza un *Relay* Teredo para crear una interfaz entre el Cliente y el *host* de destino. El *Relay* utilizado será siempre el que esté más próximo al *host* de destino, no el más próximo al cliente.

Esta técnica no es demasiado eficiente debido al *overhead* y la complejidad de su funcionamiento, pero cuando el *host* está detrás de NAT es una de las únicas opciones disponibles.

Por defecto, Windows Vista ya viene con Teredo instalado y activado, mientras que en Windows XP, 2003 y 2008 solo viene instalado. En FreeBSD y Linux ni siquiera viene instalado.

Para facilitar la comprensión del funcionamiento de este tipo de túnel, en el siguiente cuadro presentamos un resumen de los cuatro tipos de NAT existentes:

NAT de Cono - Todas las solicitudes originadas en una misma dirección y puerto internos son mapeadas a un mismo puerto de NAT. Por lo tanto solo es necesario conocer la dirección pública del NAT y el puerto asociado a un nodo interno para que un nodo externo establezca una comunicación, sin importar su dirección o puerto, pudiendo así enviar arbitrariamente paquetes al nodo interno. También se conoce como NAT Estático.

NAT de Cono Restringido - Todas las solicitudes originadas en una misma dirección y puerto internos son mapeadas a un mismo puerto de NAT. Sin embargo, el acceso externo solo se permite en respuesta a solicitudes realizadas previamente, porque la dirección del nodo externo, que está respondiendo a la solicitud, debe ser la misma utilizada como dirección de destino de la solicitud. También se conoce como NAT Dinámico.

NAT de Cono Restringido por Puerto - Tiene las mismas características de mapeo que el NAT de Cono Restringido, pero la restricción para la comunicación también considera el puerto del nodo externo. Así, un nodo externo solo podrá establecer una comunicación con un nodo interno si este último le ha enviado previamente un paquete a través del mismo puerto y dirección.

NAT Simétrico - Además de tener las mismas restricciones que el NAT tipo Cono Restringido por Puerto, cada solicitud realizada desde una dirección y puerto internos a una dirección y puerto externos es mapeada únicamente en el NAT. Es decir, si la misma dirección interna envía una solicitud, con el mismo puerto, pero para una dirección de destino diferente, en el NAT se creará un mapeo diferente. Este tipo de NAT también se conoce como NAT Masquerade o NAT Stateful.

Teredo



- Utiliza el prefijo **2001:0000::/32**.
- Los 32 bits siguientes contienen la dirección IPv4 del Servidor Teredo.
- Los 16 bits siguientes se utilizan para definir *flags* que indican el tipo de NAT utilizado e introducen una protección adicional contra los ataques de barrido.
- Los siguientes 16 bits indican el puerto UDP de salida del NAT.
- Los últimos 32 bits representan la dirección IPv4 pública del servidor NAT.

198

cgi.br

En base a los mensajes RA recibidos de los Servidores, el cliente construye su dirección IPv6 utilizando el siguiente estándar:

* Los bits 0 a 31 son el prefijo de Teredo recibido del Servidor a través de los mensajes RA; por defecto es **2001:0000**;

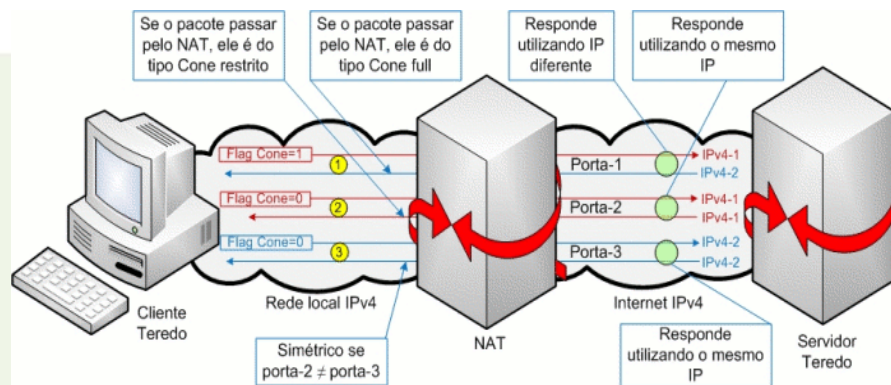
* Los bits 32 a 63 son la dirección IPv4 primaria del Servidor Teredo que envió el primer mensaje RA;

* Los bits 64 a 79 se utilizan para definir algunos *flags*. Generalmente solo se utiliza el bit más significativo, y éste se marca como 1 si el Cliente está detrás de NAT tipo Cono, caso contrario se marca como 0. Solo Windows Vista y Windows Server 2008 utilizan los 16 bits, que siguen el formato "CRAAAAUG AAAAAAAAAA", donde "C" sigue siendo el *flag* Cono; el bit R está reservado para uso futuro; el bit U define el *flag* Universal/Local (el valor por defecto es 0); el bit G define el *flag* Individual/Group (el valor por defecto 0); y los 12 bits "A" son determinados aleatoriamente por el Cliente para introducir una protección adicional contra los ataques de barrido;

* Los bits 80 a 95 son el puerto UDP de salida del NAT, recibido en los mensajes RA y enmascarado mediante la inversión de todos sus bits. Este enmascaramiento es necesario porque algunos NAT buscan puertos locales dentro de los paquetes y los reemplazan por el puerto público o viceversa;

* Los bits 96 a 127 son la dirección IPv4 pública del Servidor NAT, enmascarado a través de la inversión de todos sus bits. Este enmascaramiento es necesario porque algunos NAT buscan direcciones IP locales dentro de los paquetes y las reemplazan por la dirección pública o viceversa;

Teredo



199

cgi.br

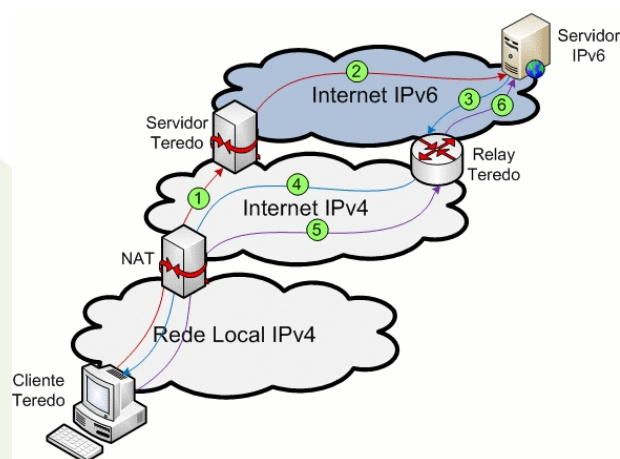
1- Se envía un mensaje *Router Solicitation* (RS) al servidor Teredo 1 (servidor primario) con el *flag* de NAT tipo Cono activado; el servidor Teredo 1 responde con un mensaje *Router Advertisement* (RA). Como el mensaje RS tenía activado el *flag* Cono, el servidor Teredo 1 envía el mensaje RA utilizando una dirección IPv4 alternativa. Con eso el cliente podrá determinar si el NAT que está utilizando es tipo Cono, si recibe el mensaje RA;

2- Si no se recibe el mensaje RA del paso anterior, el cliente Teredo envía otro mensaje RS, pero ahora con el *flag* Cono desactivado. El servidor Teredo 1 responde nuevamente con un mensaje RA pero, como el *flag* Cono del mensaje RS está desactivado, responde utilizando la misma dirección IPv4 en que recibió el mensaje RS. Si ahora el cliente recibe el mensaje RA, entonces concluye que está usando NAT tipo restringido;

3- Para tener la certeza de que el cliente Teredo no está utilizando un NAT de tipo simétrico, éste envía otro mensaje RS, esta vez al servidor secundario Teredo 2, que responde con un mensaje tipo RA. Cuando el cliente recibe el mensaje RA del servidor Teredo 2 compara la dirección y el puerto UDP de origen contenidos en el mensaje RA recibido de los dos servidores; si son diferentes el cliente concluye que está utilizando NAT tipo simétrico, que no es compatible con Teredo.

Teredo

Comunicación a través de NAT tipo CONO



200

egi.br

- 1- Para iniciar la comunicación el cliente Teredo primero debe determinar la dirección IPv4 y el puerto UDP del Relay Teredo que esté más próximo al *host* IPv6; para ello envía un mensaje ICMPv6 *echo request* al *host* IPv6 a través de su servidor Teredo;
- 2- El servidor Teredo recibe el mensaje ICMPv6 *echo request* y lo encamina al *host* IPv6 a través de la red IPv6;
- 3- El *host* IPv6 responde al cliente Teredo con un mensaje ICMPv6 *Echo Reply* que es enrutado a través del *Relay* Teredo más próximo al mismo;
- 4- Luego el *Relay* Teredo encapsula el mensaje ICMPv6 *Echo Reply* y lo envía directamente al cliente Teredo. Como el NAT utilizado por el cliente es tipo Cono, el paquete enviado por el *Relay* Teredo es encaminado al cliente Teredo;
- 5- Como el paquete que devuelve el *Relay* Teredo contiene una dirección IPv4 y el puerto UDP utilizado por el mismo, el cliente Teredo extrae esta información del paquete. Después un paquete inicial es encapsulado y enviado directamente por el cliente Teredo a la dirección IPv4 y puerto UDP del *Relay* Teredo;
- 6- El *Relay* Teredo recibe este paquete, elimina el encabezado IPv4 y UDP, y lo encamina al *host* IPv6. Luego toda la comunicación entre el cliente Teredo y el *host* IPv6 se realiza a través del *relay* Teredo con este mismo método.

Teredo

El principal problema de seguridad cuando se utiliza Teredo es que su tráfico puede pasar desapercibido por los filtros y *firewalls* si los mismos no están preparados para interpretarlo, por lo tanto las computadoras y la red interna quedan totalmente expuestos a ataques provenientes de la Internet IPv6. Para resolver este problema, antes de implementar Teredo se deben revisar los filtros y *firewalls* de la red o al menos de las computadoras que utilizarán esta técnica. Además de este problema, también existen los siguientes:

- El cliente Teredo publica en la red el puerto que abre en el NAT y el tipo de NAT que está utilizando, posibilitando así un ataque a través del mismo;
- El número de direcciones Teredo es mucho menor que las direcciones IPv6 nativas, lo que facilita la ubicación de computadoras vulnerables;
- Es fácil aplicar un ataque de denegación de servicio tanto en el cliente como en el *relay*;
- Debido al método de selección del *relay* por parte del *host* de destino, se puede crear un *relay* falso y utilizarlo para recolectar la comunicación de este *host* con sus clientes.

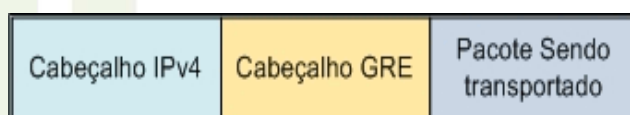
202

Más información:

- RFC 4380 - *Teredo: Tunneling IPv6 over UDP through Network Address Translations (NATs)*

GRE

- GRE (*Generic Routing Encapsulation*) - Túnel estático *host-a-host* desarrollado para encapsular diferentes tipos de protocolos.
- Soportado por la mayoría de los sistemas operativos y routers.
- Funciona tomando los paquetes originales, agregando el encabezado GRE y enviándolos a la IP de destino.
- Cuando el paquete encapsulado llega a la otra punta del túnel, el encabezado GRE se elimina y queda solo el paquete original.



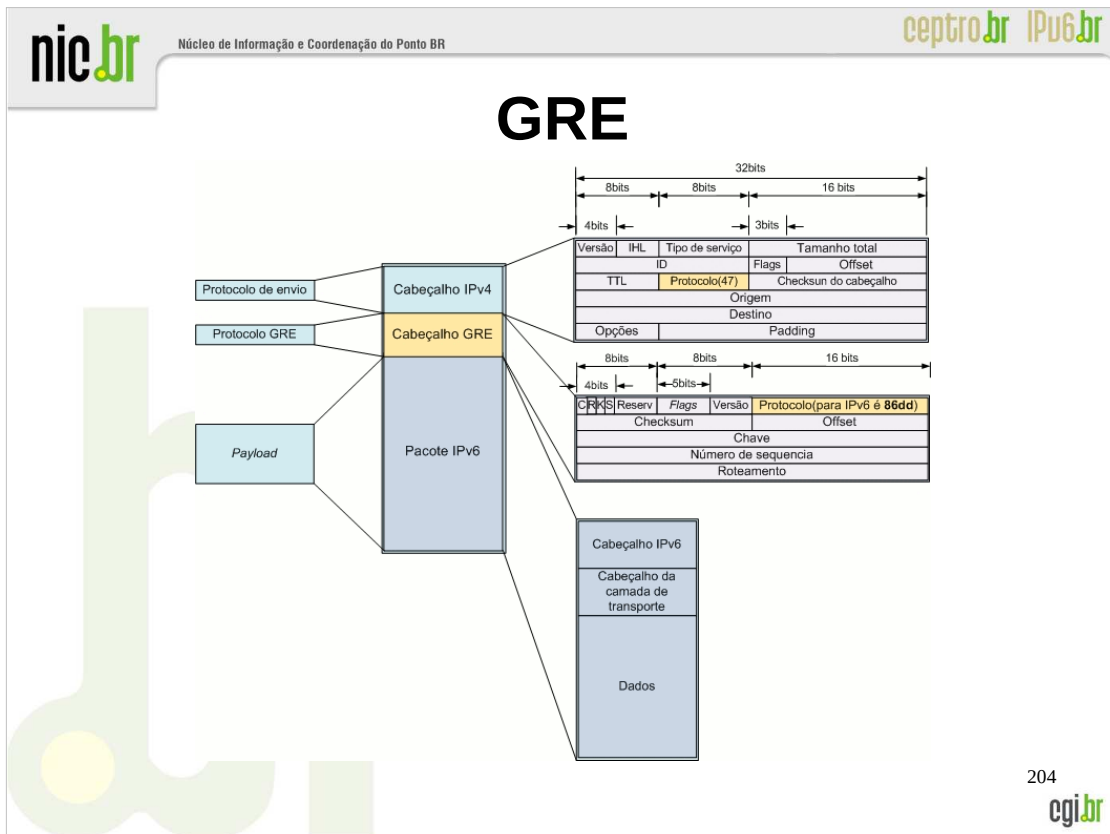
203

egi.br

GRE (*Generic Routing Encapsulation*) es un túnel estático entre dos *hosts* originalmente desarrollado por Cisco con el objetivo de encapsular diferentes tipos diferentes de protocolos, como por ejemplo IPv6 e IS-IS (ver la lista completa de protocolos soportados en <http://www.iana.org/assignments/ethernet-numbers>). Este tipo de encapsulamiento es soportado por la mayoría de los sistemas operativos y routers, y consiste en un enlace punto-a-punto. La principal desventaja del túnel GRE es la configuración manual que, dependiendo de la cantidad de túneles, implicará un gran esfuerzo en términos de su mantenimiento y administración.

Más información:

- RFC 2784 - *Generic Routing Encapsulation* (GRE)



El funcionamiento de este tipo de túnel es muy simple y consiste en tomar los paquetes originales, agregar el encabezado GRE y enviarlos a la IP de destino (la dirección de destino está especificada en el encabezado GRE); cuando el paquete encapsulado llega al otro extremo del túnel (IP de destino) se quita el encabezado GRE y queda solo el paquete original, el cual se encamina normalmente a su destinatario. Como nos preocupan más los paquetes IPv6, en el esquema siguiente mostramos la estructura de un paquete IPv6 siendo transportado en un túnel GRE:

Los campos más importantes del encabezado GRE son los siguientes:

- C (*Checksum*): Si es 1, indica que el campo *Checksum* existe y que hay información válida en el mismo y en el campo *Offset*;
- R (*Routing*): Si es 1, indica que el campo Enrutamiento existe y que hay información de enrutamiento válida en el mismo y en el campo *Offset*;
- K (*Key*): Si es 1, indica que el campo Clave existe y está siendo utilizado;
- S (*Sequence*): Si es 1, indica que el campo Número de Secuencia existe y está siendo utilizado;
- Versión Generalmente se completa con 0;
- Protocolo: Se completa con el código del protocolo que está siendo transportado, de acuerdo con los tipos de paquetes ethernet (<http://www.iana.org/assignments/ethernet-numbers>);
- *Offset*: Indica la posición donde inicia el campo de enrutamiento;
- *Checksum*: Contiene el *checksum* IP (complemento a 1) del encabezado GRE y del paquete que está siendo transportado;
- Clave (*Key*): Contiene un número de 32 bits que es introducido por el encapsulador. Es utilizado por el destinatario para identificar el remitente del paquete;
- Número de secuencia (*Sequence number*): Contiene un número entero de 32 bits que es introducido por el remitente del paquete. Es utilizado por el destinatario para secuenciar los paquetes recibidos;
- Enrutamiento (*Routing*): Contiene una lista de entradas de enrutamiento, pero generalmente no se utiliza.

Técnicas de tradução

- Posibilitan un enrutamiento transparente en la comunicación entre los nodos de una red IPv6 y los nodos de una red IPv4 y viceversa.
- Pueden actuar de diversas maneras y en capas distintas:
 - Traduciendo encabezados IPv4 en encabezados IPv6 y viceversa;
 - Convirtiendo direcciones;
 - Convirtiendo APIs de programación;
 - Actuando en el intercambio de tráfico TCP o UDP.

205

Las técnicas de traducción permiten un enrutamiento transparente en la comunicación entre nodos que solamente soportan una versión del protocolo IP o que utilizan doble pila. Estos mecanismos pueden actuar de diferentes formas y en distintas capas, traduciendo encabezados IPv4 a encabezados IPv6 y viceversa, realizando conversiones de direcciones, de APIs de programación, o actuando en el intercambio de tráfico TCP o UDP.

Más información:

- RFC 4966 - *Reasons to Move the Network Address Translator - Protocol Translator (NAT-PT) to Historic Status*

SIIT

- SIIT (*Stateless IP/ICMP Translation*) - Permite la comunicación entre nodos que solo soportan IPv6 y nodos que solo soportan IPv4.
- Utiliza un traductor ubicado en la capa de red de la pila, el cual convierte campos específicos de los encabezados de paquetes IPv6 en encabezados de paquetes IPv4 y viceversa.
- Los encabezados TCP y UDP generalmente no se traducen.
- Utiliza una dirección IPv4-mapeada en IPv6, en el formato **0::FFFF:a.b.c.d**, que identifica el destino IPv4, y una dirección IPv4-traducida, en el formato **0::FFFF:0:a.b.c.d**, para identificar el nodo IPv6.
- Utiliza rangos de direcciones IPv4 para identificar nodos IPv6.
- Traduce mensajes ICMPv4 en ICMPv6 y viceversa.

206

egi.br

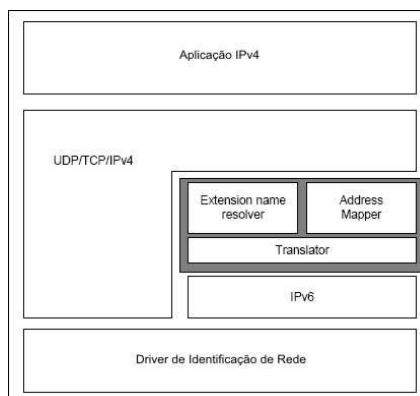
SIIT (*Stateless IP/ICMP Translation Algorithm*) - Definido en la RFC 2765, SIIT es un mecanismo de traducción *stateless* de encabezados IP/ICMP que permite la comunicación entre nodos que solo soportan IPv6 y nodos que solo soportan IPv4. Utiliza un traductor ubicado en la capa de red de la pila, el cual convierte campos específicos de los encabezados de paquetes IPv6 en encabezados de paquetes IPv4 y viceversa. Para realizar este proceso el traductor necesita una dirección IPv4-mapeada en IPv6, en formato **0::FFFF:a.b.c.d**, que identifica el destino IPv4, y una dirección IPv4-traducida, en formato **0::FFFF:0:a.b.c.d**, para identificar el nodo IPv6. Cuando el paquete llega al SIIT, el encabezado se traduce, la dirección se convierte a IPv4 y se encamina al nodo de destino;

Más información:

- RFC 2765 - *Stateless IP/ICMP Translation Algorithm (SIIT)*

BIS

- BIS (*Bump-in-the-Stack*) - Funciona entre la capa de aplicación y la capa de red.
- Se utiliza para soportar aplicaciones IPv4 en redes IPv6.
- Agrega tres módulos a la pila IPv4:
 - *Translator* - Traduce encabezados IPv4 en encabezados IPv6 y viceversa;
 - *Address mapper* - Tiene un rango de direcciones IPv4 que se asocian a direcciones IPv6 cuando el *translator* recibe un paquete IPv6;
 - *Extension name resolver* - Actúa en las consultas DNS realizadas por la aplicación IPv4.
- No funciona en comunicaciones *multicast*.



207

BIS (*Bump in the Stack*) - Este método permite la comunicación de aplicaciones IPv4 con nodos IPv6. Definida en la RFC 2767, BIS funciona entre la capa de aplicación y la capa de red, agregando a la pila IPv4 tres módulos: *translator*, que traduce los encabezados IPv4 enviados a encabezados IPv6 y los encabezados IPv6 recibidos a encabezados IPv4; *extension name resolver*, que actúa en las DNS *queries* realizadas por IPv4, de modo que, si el servidor de DNS devuelve un registro AAAA, el resolver pide al *address mapper* que asigne una dirección IPv4 correspondiente a la dirección IPv6; y *address mapper*, que posee una cierta cantidad de direcciones IPv4 para asociar a direcciones IPv6 cuando el *translator* recibe un paquete IPv6. Como las direcciones IPv4 no se transmiten en la red, éstas pueden ser direcciones privadas. Este solo método permite la comunicación de aplicaciones IPv4 con *hosts* IPv6, y no lo contrario, y además no funciona en comunicaciones *multicast*.

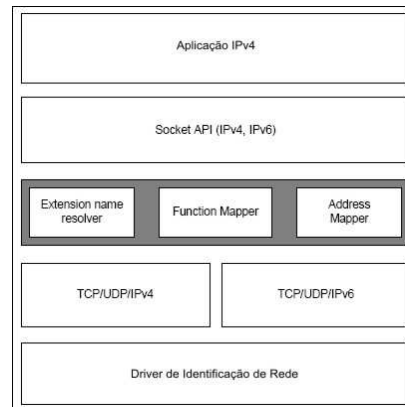
Como las direcciones IPv4 no se transmiten en la red, éstas pueden ser direcciones privadas. Este solo método permite la comunicación de aplicaciones IPv4 con *hosts* IPv6, y no lo contrario, y además no funciona en comunicaciones *multicast*.

Más información:

- RFC 2767 - *Dual Stack Hosts using the "Bump-In-the-Stack" Technique (BIS)*

BIA

- BIA (*Bump in the API*) - Similar a BIS, traduce funciones de la API IPv4 a funciones de la API IPv6 y viceversa.
- Se agregan tres módulos, **extension name resolver** y **address mapper**, que funcionan de la misma manera que en BIS, y **function mapper**, que detecta las llamadas de las funciones del socket IPv4 e invoca las funciones correspondientes del socket IPv6 y viceversa.
- También utiliza rangos de direcciones IPv4.
- No soporta comunicaciones *multicast*;



208

BIA (*Bump in the API*) - Similar al BIS, este mecanismo agrega una API de traducción entre el socket API y los módulos TPC/IP de los *hosts* de doble pila, permitiendo la comunicación entre aplicaciones IPv4 y *hosts* IPv6, traduciendo las funciones del socket IPv4 en funciones del socket IPv6 y viceversa. De acuerdo con lo descrito en la RFC 3338, se agregaron tres módulos, *extension name resolver* y *address mapper*, que funcionan de la misma manera que en BIS, y *function mapper*, que detecta las llamadas de las funciones del socket IPv4 e invoca las funciones correspondientes del socket IPv6 y viceversa. BIA tiene dos ventajas respecto a BIS: no depende del *driver* de la interfaz de red y no introduce *overhead* en la traducción de los encabezados de los paquetes. Sin embargo, tampoco soporta comunicaciones *multicast*.

Más información:

- RFC 3338 - *Dual Stack Hosts Using "Bump-in-the-API (BIA)*

TRT

- TRT (*Transport Relay Translator*) - Actúa como traductor de capa de transporte, permitiendo la comunicación entre *hosts* IPv6 e IPv4 a través de tráfico TCP/UDP.
- Actúa en equipos con doble pila que deben ser insertados en un punto intermedio de la red.
- En la comunicación de un *host* IPv6 con un *host* IPv4, agrega un prefijo IPv6 falso a la dirección IPv4 del destino.
- Cuando un paquete con ese prefijo falso atraviesa el TRT, el paquete es interceptado y enviado al *host* IPv4 de destino en un paquete TCP o UDP.
- Para que funcione de forma bidireccional se debe agregar un bloque de direcciones IPv4 públicas y usar un servidor DNS-ALG para mapear las direcciones IPv4 a IPv6.

209

TRT (*Transport Relay Translator*) - Actuando como un traductor de capa de transporte, este mecanismo permite la comunicación entre *hosts* solo IPv6 y *hosts* solo IPv4 a través de tráfico TCP/UDP. Sin necesidad de instalar ningún tipo de *software*, TRT corre en equipos con doble pila que se deben insertar en un punto intermedio dentro de la red. En la comunicación de un *host* IPv6 con un *host* IPv4, de acuerdo con la definición de la RFC 3142, a la dirección IPv4 de destino se le agrega un prefijo IPv6 falso. Cuando un paquete con este falso prefijo atraviesa el TRT, dicho paquete es interceptado y enviado al *host* IPv4 de destino en un paquete TCP o UDP. En la traducción TCP y UDP el *checksum* solo se debe recalcular en el caso de las conexiones TCP, el estado del *socket* sobre el cual el *host* está conectado debe ser mantenido, retirándolo cuando la comunicación haya finalizado. Para que el mecanismo funcione de forma bidireccional es necesario agregar un bloque de direcciones IPv4 públicas y usar un servidor DNS-ALG para mapear las direcciones IPv4 a IPv6.

Más información:

- RFC 3142 - *An IPv6-to-IPv4 Transport Relay Translator*

ALG y DNS-ALG

- ALG (*Application Layer Gateway*) - Trabaja como un *proxy* HTTP.
- El cliente inicia la conexión con el ALG, que establece una conexión con el servidor, retransmitiendo las solicitudes de salida y los datos de entrada.
- En redes solo IPv6, el ALG habilita la comunicación de los *hosts* con servicios en redes solo IPv4, configurando el ALG en nodos con doble pila.
- Normalmente se utiliza cuando el *host* que desea acceder a la aplicación en el servidor IPv4 está detrás de NAT o de un *firewall*.
- DNS-ALG - Traduce consultas DNS de tipo AAAA provenientes de un *host* IPv6 a consultas tipo A, en caso que el servidor de nombres a ser consultado se encuentre en el entorno IPv4, y viceversa.

210

Seguridad

- Con el uso de doble pila las aplicaciones quedan expuestas a ataques a ambos protocolos, IPv6 e IPv4, lo cual se puede resolver configurando *firewalls* específicos para cada protocolo.
- Las técnicas de Túneles y Traducción son las que provocan mayor impacto desde el punto de vista de la seguridad.
- Los mecanismos de tunelización son susceptibles a ataques de DoS y a la falsificación de paquetes y direcciones de los routers y *relays* utilizados por estas técnicas, como 6to4 y TEREDO.
- Las técnicas de traducción generan problemas relacionados con la incompatibilidad de dichas técnicas con algunos mecanismos de seguridad existentes. Similar a lo que ocurre con el NAT en IPv4.

Seguridad

- Como protegerse:
 - Utilizar doble pila en la migración, protegiendo ambas pilas con *firewalls*;
 - Dar preferencia a los túneles estáticos, no a los automáticos;
 - Permitir la entrada de tráfico solamente proveniente de túneles autorizados.

IPv6.br

La nueva generación del
Protocolo de Internet

213

egi.br

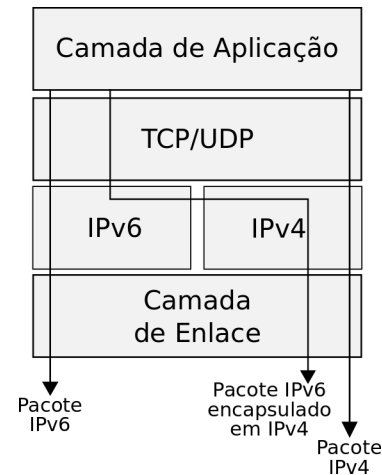
Enrutamiento IPv6

Módulo 8

En este módulo presentaremos algunas características básicas sobre el funcionamiento de los mecanismos de enrutamiento, tanto interno (IGP) como externo (EGP), siempre destacando los principales cambios en relación con IPv6. Hablaremos sobre los protocolos de enrutamiento RIP, OSPF, IS-IS y BGP.

Consideraciones importantes

- IPv4 e IPv6 → Capa de Red
- Dos redes diferentes
 - Planificación
 - Soporte
 - *Troubleshooting*
 - Arquitectura de los equipos
 - ...



215

IPv4 e IPv6 son protocolos de Capa de Red, de manera que ésta es la única capa directamente afectada por la implementación de IPv6, sin necesidad de modificar el funcionamiento de las demás.

Sin embargo, es necesario comprender que se trata de dos Capas de Red distintas e independientes. Esto implica algunas consideraciones importantes:

- Cómo actuar en la planificación y estructuración de las redes:
 - Migrar toda la estructura a doble pila; migrar solo las áreas críticas; mantener dos estructuras diferentes, una IPv4 y otra IPv6; etc.
 - En las redes con doble pila algunas configuraciones deben ser duplicadas, como por ejemplo la del DNS, el *firewall* y los protocolos de enrutamiento.
- Para el soporte y resolución de problemas será necesario si las fallas se encuentran en la conexión de la red IPv4 o de la red IPv6;
- Los nuevos equipos y aplicaciones deben soportar las funcionalidades de los dos protocolos.

Consideraciones importantes

Características fundamentales de las direcciones IP

- Identificación
 - Unívoca
 - Comandos: host, nslookup, dig...
- Localización
 - Enrutamiento y encaminamiento entre el origen y el destino
 - Comandos: mtr -4/-6, traceroute(6), tracert(6)...

Semántica sobrecargada

- Dificultad para la movilidad
- Desagregación de rutas

216

La Capa de Red está asociada principalmente a dos características:

- **Identificación** – Debe garantizar que cada dispositivo de la red sea identificado de manera unívoca, sin posibilidad de error. En otras palabras, la dirección IP debe ser única en a nivel mundial. Utilizando el comando host en las plataformas UNIX, o nslookup en las plataformas Windows, se puede verificar la identificación de un servicio, por ejemplo. En las redes con doble pila los nodos se identifican por las dos direcciones.
- **Localización** – Indica cómo llegar al destino, decidiendo el encaminamiento de los paquetes en base al direccionamiento; ocurre de la misma manera tanto en IPv4 como en IPv6. Podemos verificar esta funcionalidad utilizando comandos como *mtr -4* y *-6*, o *traceroute* (*traceroute6*), o *tracert* (*tracert6*). Estos comandos muestran la identificación y la localización de un nodo.

La unión de estas dos características en la Capa de Red resulta en una semántica sobrecargada. Esto se evidencia en aspectos tales como la agregación de rutas, agravando el problema del crecimiento de la tabla de enrutamiento global. Una forma de impedir esto consiste en separar las funciones de localización e identificación.

Consideraciones importantes

Separar las funciones de localización e identificación.

- LISP (*Locator/Identifier Separation Protocol*).
- Permite una implementación gradual.
 - No exige ningún cambio en las pilas de los *host* ni grandes cambios en la infraestructura existente.
- EID (*Endpoint Identifiers*).
- RLOC (*Routing Locators*).
- ITR (*Ingress Tunnel Router*) / ETR (*Egress Tunnel Router*).
 - Realiza el mapeo entre EID y RLOC.
- Utiliza tanto IPv4 como IPv6.

217

Hay un grupo de trabajo en la IETF que está discutiendo una forma de separar esas dos funciones (identificación y localización). LISP (*Locator/Identifier Separation Protocol*) es un protocolo sencillo que busca separar las direcciones IP en *Endpoint Identifiers* (EIDs) y *Routing Locators* (RLOCs). No exige ningún cambio en las pilas de los *host* ni grandes cambios en la infraestructura existente, pudiendo ser implementado en un número relativamente pequeño de routers.

Sus principales elementos son los siguientes:

- *Endpoint ID* (EID): Un identificador de 32 bits (para IPv4) o 128 bits (para IPv6) que se utiliza en los campos de dirección de origen y destino del primer encabezado (más interno) de un paquete. El *host* obtiene un EID de destino de la misma manera en que hoy obtiene una dirección de destino, por ejemplo a través de una consulta de DNS. El EID de origen también se obtiene a través de los mecanismos ya existentes usados para definir la dirección local de un *host*;
- *Routing Locator* (RLOC): Dirección IPv4 o IPv6 de un ETR (*Egress Tunnel Router*). Los RLOC se numeran a partir de un bloque topológicamente agregado, y son atribuidos a una red en cada punto en que hay conexión a la Internet global;
- *Ingress Tunnel Router* (ITR): Router de entrada del túnel que recibe un paquete IP (más precisamente, un paquete IP que no contiene un encabezado LISP), trata la dirección de destino de ese paquete como un EID y realiza un mapeo entre el EID y el RLOC. A continuación el ITR agrega un encabezado "IP externa" que contiene uno de sus RLOC globalmente ruteables en el campo de dirección de origen, y un RLOC – resultado del mapeo – en el campo de dirección de destino;
- *Egress Tunnel Router* (ETR): Router de salida del túnel que recibe un paquete IP donde la dirección de destino del encabezado "IP externa" es uno de sus RLOC. El router retira el encabezado externo y encamina el paquete en base al siguiente encabezado IP encontrado.

Más información

- *Locator/ID Separation Protocol (LISP)* - <http://www.ietf.org/id/draft-ietf-lisp-06.txt>
- *LISP Networking: Topology, Tools, and Documents* - <http://www.lisp4.net> (solo conexiones IPv4)
- *LISP Networking: Topology, Tools, and Documents* - <http://www.lisp6.net> (solo conexiones IPv6)

Consideraciones importantes

Prefijo IP

- El recurso que el Registro.br asigna al AS es un bloque IP.
- El bloque IP no es ruteable.
 - El bloque es un grupo de IPs.
- El prefijo IP es ruteable.
 - Número de bits que identifica la red;
 - Puede crear un prefijo /32 igual al bloque /32 IPv6 recibido del Registro.br;
 - Puede crear un prefijo /33, /34,... /48.
- Esta nomenclatura es importante.
 - Activación de sesiones de tránsito con otros operadores;
 - *Troubleshooting*.

218

La definición de prefijo IP es una definición importante.

El recurso que asigna el Registro.br a los AS es un bloque IP, el cual representa un grupo de direcciones IP. Un bloque no es un elemento ruteable; lo que es ruteable es el prefijo. Lo que es posible, por ejemplo, es crear un prefijo IPv6 /32 igual al bloque /32 recibido del Registro.br y anunciar dicho prefijo en la tabla de rutas. Pero a partir del bloque recibido también se pueden crear prefijos /33, /34, /48 etc.

El prefijo representa el número de bits de una dirección que identifica la red.

A pesar de ser solo una nomenclatura, esta definición es importante a la hora de enviar información para activar sesiones de tránsito con otros operadores y en la detección de problemas de conectividad.

¿Cómo funciona el router?

Ejemplo:

1. El router recibe una trama Ethernet;
2. Verifica la información del Ethertype que indica que el protocolo de capa superior transportado es IPv6;
3. Se procesa el encabezado IPv6 y se analiza la dirección de destino;
4. El router busca en la tabla de enrutamiento *unicast* (RIB - *Router Information Base*) si hay alguna entrada a la red de destino;

- Visualización de la RIB:

show ip(v6) route → Cisco/Quagga

show route (table inet6) → Juniper

219

cgi.br

También es importante comprender el funcionamiento básico de un router y de qué manera procesa los paquetes recibidos y toma las decisiones de encaminamiento. Consideremos el siguiente ejemplo:

- El router recibe una trama Ethernet a través de su interfaz de red;
- Verifica la información del Ethertype que indica que el protocolo de capa superior transportado es IPv6;
- Se procesa el encabezado IPv6 y se analiza la dirección de destino;
- El router busca en la tabla de enrutamiento *unicast* (RIB - *Router Information Base*) si hay alguna entrada a la red de destino;
-

Visualización de la RIB IPv6:

Cisco/Quagga → show ipv6 route

Juniper → show route table inet6

Visualización de la RIB IPv4:

Cisco/Quagga → show ip route

Juniper → show route

¿Cómo funciona el router?

5. *Longest Match* - Busca la entrada más específica. Ejemplo:

- La IP de destino es 2001:0DB8:0010:0010::0010
- El router tiene la siguiente información en su tabla de rutas:
 - 2001:DB8::/32 vía interfaz A
 - 2001:DB8::/40 vía interfaz B
 - 2001:DB8:10::/48 vía interfaz C
- Los tres prefijos engloban la dirección de destino, pero el router siempre preferirá el más específico, en este caso el /48;
- ¿Cuál es la entrada IPv4 e IPv6 más específica?

6. Una vez identificado el prefijo más específico, el router decrementa el *Hop-Limit*, arma la trama Ethernet de acuerdo con la interfaz y envía el paquete.

220

- *Longest Match* - Busca la entrada más específica. Ejemplo:
 - La IP de destino es 2001:0DB8:0010:0010::0010
 - El router tiene la siguiente información en su tabla de rutas:
 - 2001:DB8::/32 vía interfaz A
 - 2001:DB8::/40 vía interfaz B
 - 2001:DB8:10::/48 vía interfaz C
 - Los tres prefijos engloban la dirección de destino, pero el router siempre preferirá el más específico, en este caso el /48;
 - Una vez identificado el prefijo más específico, el router decrementa el *Hop-Limit*, arma la trama Ethernet de acuerdo con la interfaz y envía el paquete.

¿Cómo funciona el router?

¿Qué pasa si hay más de un camino para el mismo prefijo?

- Se utiliza una tabla de preferencias predefinida.
- Número entero comprendido entre 0 y 255 asociado a cada ruta; cuanto menor sea su valor más confiable será la ruta;
- Evalúa si está conectado directamente, si la ruta fue aprendida a través del protocolo de enrutamiento externo o interno;
- Tiene significado local, no puede ser anunciado por los protocolos de enrutamiento;
- Su valor puede ser modificado en caso que sea necesario priorizar un determinado protocolo.

¿Qué pasa si el valor de la tabla de preferencias también es el mismo?

221

cgi.br

Si el router encuentra más de un camino para el mismo destino con el mismo valor de *longest match*, éste una tabla de preferencias predefinida (concepto de *Distancia Administrativa* de Cisco).

Los valores de esta tabla son números enteros comprendidos entre 0 y 255 asociados a cada ruta; cuanto menor es su valor más confiable es la ruta; Los valores se asignan evaluando si la ruta está conectada directamente, si fue aprendida a través del protocolo de enrutamiento externo o interno, etc. Estos valores solo tienen significado local, no pueden ser anunciados por los protocolos de enrutamiento y, si fuera necesario, pueden ser modificados para priorizar un determinado protocolo.

En caso que en la tabla de preferencias también se encuentre el mismo valor, hay equipos e implementaciones que por defecto realizan el balanceo de carga.

Tabla de Enrutamiento

- El proceso de selección de rutas es idéntico en IPv4 e IPv6, pero las tablas de rutas son independientes.
 - Hay una RIB IPv4 y otra IPv6.
- A través de mecanismos de optimización, las mejores rutas se agregan a la tabla de encaminamiento
 - FIB - *Forwarding Information Base*;
 - La FIB se crea a partir de la RIB;
 - Al igual que la RIB, la FIB también está duplicada.
- En los routers que tienen arquitectura distribuida el proceso de selección de rutas y el encaminamiento de los paquetes son funciones diferentes.

222

El proceso de selección de rutas es idéntico en IPv4 e IPv6, pero las tablas de rutas son independientes. Por ejemplo: Hay una RIB IPv4 y otra IPv6.

Para optimizar el envío de paquetes hay mecanismos que agregan solo las mejoras rutas a otra tabla, la tabla de encaminamiento (FIB - *Forwarding Information Base*). Un ejemplo de este mecanismo es el CEF (*Cisco Express Forwarding*) de Cisco.

La FIB se crea a partir de la RIB y, al igual que la RIB, también está duplicada si la red está configurada con doble pila. Es así que hay más información para almacenar y procesar.

En los routers que tienen arquitectura distribuida el proceso de selección de rutas y el encaminamiento de los paquetes son funciones diferentes.

Ejemplo:

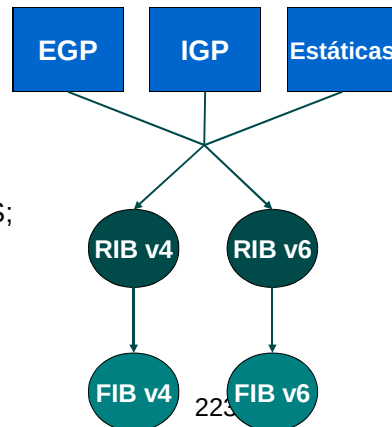
- Routers 7600 de Cisco: La RIB reside en el módulo de enrutamiento central y la FIB en las placas de las interfaces.
- Routers Juniper de la serie M: El *Router Engine* es responsable por la RIB, mientras que la FIB también reside en las placas de las interfaces (*Packet Forwarding Engine* - PFE).

Tabla de Enrutamiento

- Son las informaciones recibidas por los protocolos de enrutamiento que "alimentan" la RIB, la cual a su vez "alimenta" la FIB.

- Los Protocolos de Enrutamiento se dividen en dos grupos:

- **Interno (IGP)** - Protocolos que distribuyen la información de los routers dentro de Sistemas Autónomos. Ejemplo: OSPF; IS-IS; RIP.
- **Externo (EGP)** - Protocolos que distribuyen la información entre Sistemas Autónomos. Ejemplo: BGP-4.



Es el mecanismo de enrutamiento que permite el encaminamiento de paquetes de datos entre dos dispositivos cualquiera conectados a Internet.

Para actualizar la información que utilizan los routers para encontrar el mejor camino disponible en el encaminamiento de los paquetes hasta su destino se utilizan los protocolos de enrutamiento. Son las informaciones recibidas por los protocolos de enrutamiento que "alimentan" la RIB, la cual a su vez "alimenta" la FIB.

Estos protocolos se dividen en dos grupos:

Interno (IGP) - Protocolos que distribuyen la información de los routers dentro de Sistemas Autónomos. Como ejemplo de estos protocolos podemos mencionar: OSPF, IS-IS y RIP.

- **Externo (EGP)** - Protocolos que distribuyen la información entre Sistemas Autónomos. Como ejemplo podemos mencionar el protocolo BGP-4.

Ruta por defecto

- Cuando un router no encuentra una entrada en la tabla de rutas para una determinada dirección, ese router utiliza una ruta por defecto.
- Los servidores, estaciones de trabajo, *firewalls*, etc. solo conocen las redes directamente conectadas a una interfaz.
 - Para llegar a un destino que no esté directamente conectado deberán usar la ruta por defecto hacia otro que sí conozcan.
- ¿Todo el mundo necesita tener una ruta por defecto?

224

cgi.br

Si el router recibe un paquete cuya dirección de destino no esté explícitamente listada en la tabla de rutas, éste utilizará su ruta *por defecto*.

Naturalmente, los servidores y estaciones de trabajo necesitan una ruta por defecto. Estos dispositivos no son equipos de red; solo conocen las redes directamente conectadas a sus interfaces. Para llegar a un destino que no esté directamente conectado deberán usar la ruta por defecto hacia otro que sí conozcan.

Aquí surge la siguiente pregunta: ¿Todo el mundo necesita tener una ruta por defecto?

Ruta por defecto

- DFZ (*Default Free Zone*) - Concepto que existe entre los operadores. Es una región de Internet que no tiene ruta por defecto.
- Los routers DFZ no tienen ruta por defecto, tienen la tabla BGP completa.
- ¿Los AS que tienen la tabla completa deben tener ruta *por defecto*?
- La tabla completa muestra todas las entradas de red del mundo.
 - Los routers deben procesar información del mundo entero en tiempo real;
 - Problemas de escalabilidad futura.

225

Entre los operadores existe un concepto que delimita una región de Internet sin ruta por defecto, la DFZ (*Default Free Zone*).

Un AS que tiene la tabla completa no necesita tener ruta por defecto, ya que la tabla completa muestra las entradas d red de todo el mundo.

Este modelo es bueno y funcional, pero puede acarrear algunos problemas. Los routers deben procesar información del mundo entero en tiempo real; también pueden surgir problemas de escalabilidad futura.

Ruta por defecto

- Si hay tabla completa y ruta por defecto, ¿se utiliza la ruta por defecto?
- Ejemplo:
 - Imagine una red comprometida por un *malware*;
 - La máquina contaminada “barrerá” Internet intentando contaminar otras máquinas, incluso IPs que no están asignadas y que no están en la tabla completa;
 - Si hay ruta por defecto, su router va a encaminar ese tráfico no válido hacia adelante;
 - Este es uno de los motivos para utilizar DFZ;
 - Sugerencia: Crear una ruta por defecto y apuntar hacia Null0 o DevNull, deshabilitando el envío de mensajes '*ICMP unreachable*'.
- La ruta por defecto en IPv4 es 0.0.0.0/0 y en IPv6 ::/0.

226

El uso de ruta por defecto por parte de los routers que tienen tabla completa puede ocasionar algunos problemas.

Como ejemplo, imagina la siguiente situación: una red ha sido comprometida por un malware. La máquina contaminada “barrerá” Internet intentando contaminar otras máquinas, incluso IPs que no están asignadas y que no están en la tabla completa; Si hay ruta por defecto, su router va a encaminar ese tráfico no válido hacia adelante; Este es uno de los motivos para utilizar DFZ. Una sugerencia para solucionar este problema es crear una ruta por defecto y apuntar hacia Null0 o DevNull. También hay que deshabilitar el envío de mensajes '*ICMP unreachable*': ya que cuando un router descarta un paquete envía un mensaje '*ICMP unreachable*' pero si el destino no es válido no es necesario avisar al origen, esto solo consumirá CPU innecesariamente.

La ruta por defecto en IPv4 es 0.0.0.0/0 y en IPv6 ::/0.

Protocolos de Enrutamiento Interno

- Hay dos opciones principales para trabajar con el enrutamiento interno:
 - OSPF
 - IS-IS
 - protocolos tipo *Link-State*;
 - consideran la información de estado y envían actualizaciones de manera optimizada;
 - trabajan con estructura jerárquica.
- Tercera opción
 - RIP
- El protocolo de enrutamiento interno solo debe ser habilitado en las interfaces necesarias.

227

Hoy en día hay dos opciones principales para trabajar con enrutamiento interno, OSPF e IS-IS. Estos dos protocolos son de tipo *Link-State*, es decir, consideran la información de estado del enlace y envían actualizaciones en forma optimizada solo cuando se producen cambios de estado. También permiten trabajar con estructura jerárquica, separando la red por regiones. Esto es un punto fundamental para IPv6.

Otra opción es el protocolo RIP (*Routing Information Protocol*). Éste es un protocolo de tipo Vector de Distancia (Bellman-Ford), de fácil implementación y de funcionamiento sencillo, pero presenta algunas limitaciones como el hecho de enviar su tabla de estados periódicamente sin importar si hay o no cambios en la red.

Es importante que el protocolo de enrutamiento interno se habilite solamente en las interfaces donde sea necesario. Aunque parezca obvio, hay quienes lo configuran equivocadamente haciendo que los routers queden intentando establecer relaciones de vecindad con otros AS.

RIPng

- *Routing Information Protocol next generation* (RIPng) - Protocolo IGP simple y de fácil implementación y configuración.
- Protocolo de tipo Vector de Distancia (Bellman-Ford).
- Basado en el protocolo RIPv2 (IPv4).
- Protocolo específico para IPv6.
 - Soporte para el nuevo formato de direcciones;
 - Utiliza la dirección *multicast* **FF02::9** (*All RIP Routers*) como destino;
 - La dirección del salto siguiente debe ser una dirección *link local*;
 - En un ambiente IPv4+IPv6 es necesario utilizar RIP (IPv4) y RIPng (IPv6).

Para tratar el enrutamiento interno IPv6 se definió una nueva versión del protocolo RIP, el *Routing Information Protocol next generation* (RIPng). Esta versión se basa en el protocolo RIPv2 que se utiliza en las redes IPv4, pero es específica para redes IPv6.

Entre los principales cambios se destacan los siguientes:

- Soporte para el nuevo formato de direcciones;
- Utiliza la dirección *multicast* FF02::9 (*All RIP Routers*) como destino;
- La dirección del salto siguiente debe ser una dirección *link local*.

En un ambiente con doble pila (IPv4+IPv6) es necesario utilizar una instancia de RIP para IPv4 y una de RIPng para el enrutamiento IPv6.

A pesar de ser nuevo, el protocolo RIPng todavía tiene las mismas limitaciones que las versiones anteriores utilizadas con IPv4, entre ellas:

- El diámetro máximo de la red es de 15 saltos;
- Para determinar el mejor camino utiliza solamente la distancia;
- *Loops* de enrutamiento y conteo al infinito.

Más información:

- RFC 2080 - *RIPng for IPv6*

RIPng

- Limitaciones:
 - El diámetro máximo de la red es de 15 saltos;
 - Para determinar el mejor camino utiliza solamente la distancia;
 - *Loops* de enrutamiento y conteo al infinito.
- Actualización de la tabla de rutas:
 - Envío automático cada 30 segundos - sin importar si hay cambios o no.
 - Cuando detecta cambios en la topología de la red - envía solo la línea afectada por el cambio.
 - Cuando recibe un mensaje de tipo *Request*.

La tabla de rutas contiene la siguiente información:

- Prefijo de destino
- Métrica
- Siguiendo salto
- Identificación de la ruta (*route tag*)
- Cambio de ruta
- Tiempo hasta que la ruta expira (por defecto 180 segundos)
- Tiempo hasta el *garbage collection* (por defecto 120 segundos)

Las tablas de rutas se pueden actualizar de tres maneras: a través del envío automático de datos cada 30 segundos; cuando se detecta algún cambio en la topología de la red, enviando solo la línea afectada por el cambio; y cuando se recibe un mensaje de tipo *Request*.

RIPng

- Mensajes *Request* y *Response*

8 bits	8 bits	16 bits
Comando	Versión	Reservado
Entrada 1 de la tabla de rutas (RTE)		
....		
Entrada n de la tabla de rutas		

- RTE

- Prefijo IPv6 (128 bits)
- Identificación de la ruta (16 bits)
- Tamaño del prefijo (8 bits)
- Métrica (8 bits)
- A diferencia de lo que ocurre en RIPv2, la dirección del siguiente salto aparece solo una vez, seguida por todas las entradas que deben utilizarla.

El encabezado de los mensajes RIPng es muy sencillo y está formado por los siguientes campos:

- Comando (*command*) – Indica si el mensaje es de tipo *Request* o *Response*;
- Versión (*version*) – Indica la versión del protocolo que actualmente es 1.

Estos campos van seguidos por las entradas de la tabla de rutas (*Route Table Entry* – RTE):

- Prefijo IPv6 (128 bits);
- Identificación de la ruta (16 bits);
- Tamaño del prefijo (8 bits);
- Métrica (8 bits).

A diferencia de lo que ocurre en RIPv2, la dirección del siguiente salto aparece solo una vez, seguida por todas las entradas que deben utilizarla.

OSPFv3

- *Open Shortest Path First version 3* (OSPFv3) – Protocolo IGP de tipo *link-state*
- Los routers describen su estado actual a lo largo del AS enviando LSAs (*flooding*)
- Utiliza el algoritmo del camino más corto de Dijkstra
- Agrupa los routers en áreas
- Basado en el protocolo OSPFv2
- Protocolo específico para IPv6
- En un ambiente IPv4+IPv6 es necesario utilizar OSPFv2 (IPv4) y OSPFv3 (IPv6).

OSPF es un protocolo de tipo *link-state* donde, a través del proceso de *flooding* (inundación), los routers envían *Link State Advertisements* (LSA) describiendo su estado actual a lo largo del AS. El *flooding* consiste en el envío de un LSA por todas las interfaces de salida del router, de modo que todos los routers que reciben un LSA también lo envían por todas sus interfaces. De este modo, el conjunto de los LSAs de todos los routers forma una base de datos del estado del enlace, donde cada router que participa del AS tiene una base de datos idéntica. Con la información de esta base de datos, el router, a través del protocolo OSPF, construye un mapa de la red que será utilizado para determinar un árbol de caminos más cortos dentro de toda la subred, teniendo al propio nodo como raíz. Utiliza el algoritmo de Dijkstra para escoger el mejor camino y permite agrupar los routers en áreas.

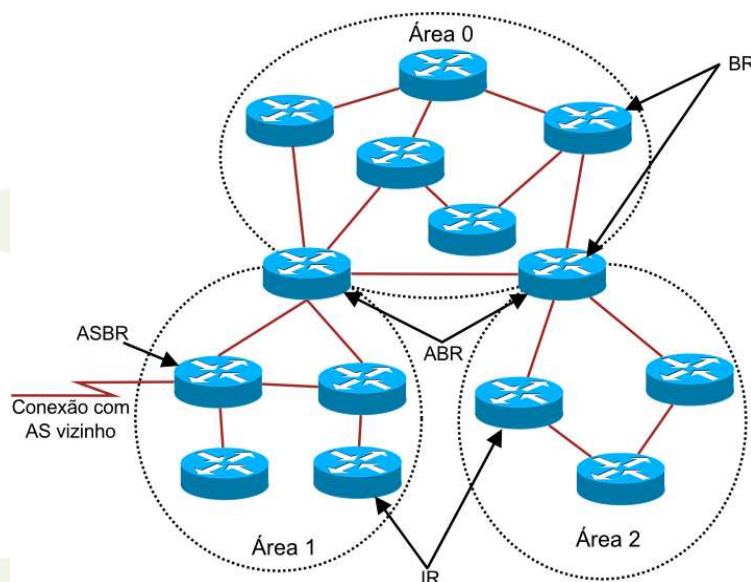
OSPF se puede configurar para trabajar de forma jerárquica, dividiendo los routers de un AS en diferentes áreas. A cada una de estas áreas se atribuye un identificador único (Area-ID) de 32 bits y todos los routers de una misma área mantienen base de datos de estado separada, de modo que la topología de una área se desconoce fuera de la misma, reduciendo así la cantidad de tráfico de enrutamiento entre las partes del AS. El área *backbone* es la responsable de distribuir la información de enrutamiento entre las áreas *nonbackbone* y se identifica mediante el ID 0 (o 0.0.0.0). En los AS en los cuales no hay esta división, generalmente el área *backbone* es la única que se configura.

A pesar de que está basado en la versión de OSPFv2 que se utiliza en las redes IPv4, OSPFv3 es un protocolo específico para IPv6. Por lo tanto, en las redes con doble pila es necesario utilizar OSPFv2 para realizar el enrutamiento IPv4 y OSPFv3 para realizar el enrutamiento IPv6.

Más información:

- RFC 5340 - *OSPF for IPv6*

Routers OSPFv3



cgi.br

Los routers OSPF se pueden clasificar de la siguiente manera:

- *Internal Router (IR)* – Routers que solo se relacionan con vecinos OSPF de una misma área;
- *Area Border Router (ABR)* – Routers que conectan una o más áreas al *backbone*. Estos poseen múltiples copias de las bases de datos de estado, una para cada área, y son responsables por condensar la información de estas áreas y enviarla al *backbone*;
- *Backbone Router (BR)* – Routers que pertenecen al área *backbone*. Un ABR es siempre un BR, ya que todas sus áreas están directamente conectadas al *backbone* o conectadas vía *virtual link* - túnel que conecta una área al *backbone* pasando a través de otra área; y
- *Autonomous System Border Router (ASBR)* – Routers que intercambian información de enrutamiento con routers de otro AS y distribuyen las rutas recibidas dentro su propio AS.

OSPFv3

Semejanzas entre OSPFv2 y OSPFv3

- Tipos básicos de paquetes
 - Hello, DBD, LSR, LSU, LSA
- Mecanismos para descubrimiento de vecinos y formación de adyacencias
- Tipos de interfaces
 - *point-to-point*, *broadcast*, NBMA, *point-to-multipoint* y enlaces virtuales
- Lista de estados y eventos de las interfaces
- Algoritmo de selección del *Designated Router* y del *Backup Designated Router*
- Envío y edad de las LSAs
- AREA_ID y ROUTER_ID continúan siendo de 32 bits

OSPFv3 todavía incluye algunas características de OSPFv2:

- Tipos básicos de paquetes
 - Hello, DBD, LSR, LSU, LSA
- Mecanismos para descubrimiento de vecinos y formación de adyacencias
- Tipos de interfaces
 - *point-to-point*, *broadcast*, NBMA, *point-to-multipoint* y enlaces virtuales
- Lista de estados y eventos de las interfaces
- Algoritmo de selección del *Designated Router* y del *Backup Designated Router*
- Envío y edad de las LSAs
- AREA_ID y ROUTER_ID continúan siendo de 32 bits

OSPFv3

Diferencias entre OSPFv2 y OSPFv3

- OSPFv3 funciona por enlace, y no por subred
- Se eliminó la información de direccionamiento
- Se agregó limitación de alcance para *flooding*
- Soporte explícito para múltiples instancias en cada enlace
- Uso de direcciones *link-local*
- Cambios en la autenticación
- Cambios en el formato del paquete
- Cambios en el formato del encabezado LSA
- Tratamiento de tipos de LSA desconocidos
- Soporte para áreas Stub/NSSA
- Identificación de vecinos mediante Router IDs
- Utiliza direcciones *multicast* (*AllSPFRouters* **FF02::5** y *AllDRouters* **FF02::6**)

Entre las principales diferencias entre OSPFv2 y OSPFv3 se destacan las siguientes:

- OSPFv3 funciona por enlace, y no por subred
- Se eliminó la información de direccionamiento
- Se agregó limitación de alcance para *flooding*
- Soporte explícito para múltiples instancias en cada enlace
- Uso de direcciones *link-local*
- Cambios en la autenticación
- Cambios en el formato del paquete
- Cambios en el formato del encabezado LSA
- Tratamiento de tipos de LSA desconocidos
- Soporte para áreas Stub/NSSA
- Identificación de vecinos mediante Router IDs
- Utiliza direcciones *multicast* (*AllSPFRouters* **FF02::5** y *AllDRouters* **FF02::6**)

IS-IS

- *Intermediate System to Intermediate System* (IS-IS) - Protocolo IGP de tipo *link-state*
- Originalmente desarrollado para funcionar sobre el protocolo CLNS
 - *Integrated IS-IS* permite enrutar tanto IP como OSI
 - Utiliza NLPID para identificar el protocolo de red utilizado
- Trabaja en dos niveles
 - L2 = Backbone
 - L1 = Stub
 - L2/L1= Interconexión L2 y L1

Al igual que OSPF, *Intermediate System to Intermediate System* (IS-IS) es un protocolo IGP de tipo *link-state*, que utiliza el algoritmo de Dijkstra para calcular las rutas.

IS-IS fue originalmente desarrollado para funcionar sobre el protocolo CLNS, pero la versión *Integrated IS-IS* permite enrutar tanto paquetes de red IP como OSI. Para ello se utiliza un identificador de protocolo, el NLPID, para informar qué protocolo de red está siendo utilizado.

Al igual que OSPF, IS-IS también permite trabajar la red de manera jerárquica, actuando con los routers en dos niveles, L1 (Stub) y L2 (Backbone), además de los routers que integran esas áreas, L2/L1.

IS-IS

- No se ha desarrollado una nueva versión para trabajar con IPv6. Solo se han agregado nuevas funcionalidades a la versión ya existente
- Dos nuevos TLVs para
 - *IPv6 Reachability*
 - *IPv6 Interface Address*
- Nuevo identificador de capa de red
 - IPv6 NLPID
- El proceso de establecimiento de vecindades no cambia

Para tratar el enrutamiento IPv6 no se definió una nueva versión del IS-IS sino que solo se agregaron nuevas funcionalidades a la versión ya existente.

Se agregaron dos nuevas TLVs (*Type-Length-Values*):

- **IPv6 Reachability** (type 236) – Transporta información de las redes accesibles;
- **IPv6 Interface Address** (type 232) – Indica las direcciones IP de la interfaz que está transmitiendo el paquete.

También se agregó un nuevo identificador de la capa de red

- **IPv6 NLPID** – Su valor es 142.

El proceso de establecimiento de vecindades no cambia.

Más información:

- RFC 1195 - *Use of OSI IS-IS for Routing in TCP/IP and Dual Environments*
- RFC 5308 - *Routing IPv6 with IS-IS*

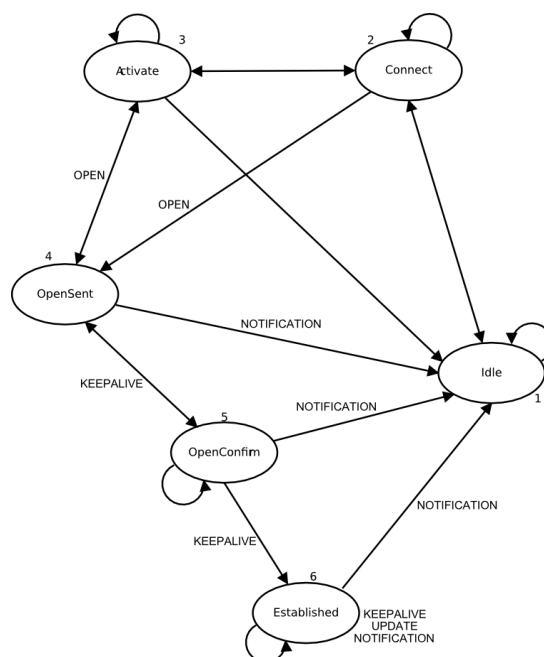
Protocolo de Enrutamiento Externo

- En la actualidad el protocolo de enrutamiento externo por defecto es *Border Gateway Protocol* versión 4 (BGP-4).
 - protocolo de tipo *path vector*.
- Los routers BGP intercambian información de enrutamiento entre ASs vecinos.
 - con esta información diseñan un grafo de conectividad entre los AS.

En la actualidad el protocolo de enrutamiento externo por defecto es *Border Gateway Protocol* versión 4 (BGP-4). Se trata de un protocolo de tipo *path vector*, en el cual los routers BGP intercambian información de enrutamiento entre ASs vecinos diseñando un grafo de conectividad entre los mismo.

BGP

- Puerto TCP 179
- Cuatro tipos de mensajes:
 - *Open*
 - *Update*
 - *Keepalive*
 - *Notification*
- Dos tipos de conexión:
 - eBGP
 - iBGP
- Funcionamiento representado por una Máquina de Estados.



BGP es un protocolo extremadamente simple basado en sesiones TCP escuchando en el puerto 179.

Para intercambiar información y mantener el estado de la conexión TCP se utilizan cuatro tipos de mensajes BGP:

- *Open* – Enviado por los dos vecinos luego del establecimiento de la conexión TCP, lleva la información necesaria para el establecimiento de la sesión BGP (ASN, versión de BGP, etc);
- *Update* – Usado para transferir la información de enrutamiento entre los vecinos BGP, la cual se utilizará para construir el grafo que describe la relación entre varios ASs;
- *Keepalive* – Se envían frecuentemente para evitar que la conexión TCP expire;
- *Notification* – Se envía cuando se detecta un error, cerrando la conexión BGP inmediatamente después de su envío.

Usted puede establecer dos tipos de conexión BGP:

- externa (eBGP) – conexión entre dos AS vecinos;
- interna (iBGP) – conexión entre routers dentro de un mismo AS. Establecer el iBGP es muy importante para mantener una visión consistente de las rutas externas en todos los routers de un AS.

El funcionamiento de BGP se puede representar mediante una Máquina de Estados Finitos. Para quien no están familiarizado con el protocolo BGP, al verificar que el estado de una conexión está “Active” o “Established”, puede tener la falsa impresión de que la conexión está “activa” o “establecida”, pero en general, en BGP, cuando hay “palabras” representando el estado, significa que la sesión BGP todavía no está bien. La sesión estará efectivamente establecida cuando se observe el número de prefijos que se está recibiendo del vecino. Esos nombres representan estados intermedios de la sesión BGP. Identificar esos estados ayuda en el análisis y resolución de problemas.

Más información:

RFC 4271 - *A Border Gateway Protocol 4 (BGP-4)*

RFC 4760 - *Multiprotocol Extensions for BGP-4*

Atributos BGP

- El criterio de selección entre diferentes atributos BGP varía de implementación a implementación.
- Los atributos BGP se dividen en categorías y subcategorías.

<i>ORIGIN</i>	Bien conocido	Obligatorio
<i>AS_PATH</i>	Bien conocido	Obligatorio
<i>NEXT_HOP</i>	Bien conocido	Obligatorio
<i>MULTI_EXIT_DISC</i>	Opcional	No transitivo
<i>LOCAL_PREF</i>	Bien conocido	Discrecional
<i>ATOMIC_AGGREGATE</i>	Bien conocido	Discrecional
<i>AGGREGATOR</i>	Opcional	Transitivo

A pesar de que la RFC sobre BGP recomienda algunos puntos, el criterio de selección entre diferentes atributos BGP puede variar de implementación a implementación. Sin embargo, la mayor parte de las implementaciones sigue los mismos estándares.

Los atributos BGP se pueden dividir en dos grandes categorías:

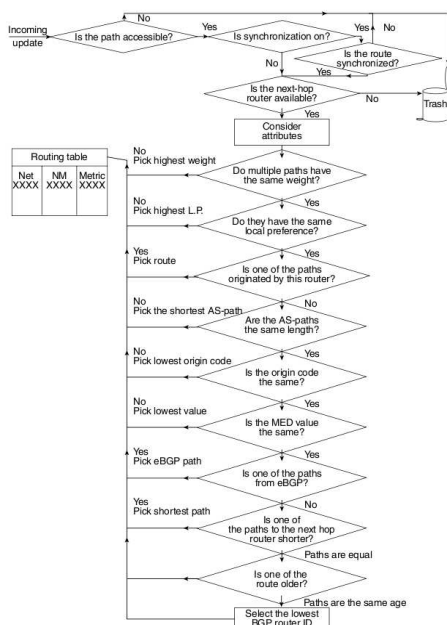
- **Bien conocidos** (*Well-known*) – Son atributos definidos en la especificación original del protocolo BGP. Estos se subdividen en otras dos categorías:
 - **Obligatorios** (*Mandatory*) – Siempre deben estar presentes en los mensajes tipo UPDATE y deben ser obligatoriamente reconocidos por todas las implementaciones del protocolo;
 - **Discrecional** (*Discretionary*) – No deben obligatoriamente estar presentes en todos los mensajes UPDATE.
- **Opcionales** (*Optional*) – No son obligatoriamente soportados por todas las implementaciones de BGP. Estos se subdividen en otras dos categorías:
 - **Transitivos** (*Transitive*) – Deben ser re-transmitidos en los mensajes UPDATE;
 - **No Transitivos** (*Non-transitive*) – No deben ser re-transmitidos.

La RFC sobre BGP contiene los siguientes atributos:

- *ORIGIN* – Es bien conocido y obligatorio. Indica si el camino fue aprendido vía IGP o EGP;
- *AS_PATH* - Es bien conocido y obligatorio. Indica el camino para llegar a un destino, listando los ASN por los cuales se debe pasar;
- *NEXT_HOP* – Es bien conocido y obligatorio. Indica la dirección IP de la interfaz del siguiente router;
- *MULTI_EXIT_DISC* – Es opcional y no transitivo. Indica a los vecinos BGP externos cuál es el mejor camino para una determinada ruta del propio AS, influenciándolos, así como cuál camino se debe seguir en caso que el AS posea diferentes puntos de entrada;
- *LOCAL_PREF* – Es bien conocido y discrecional. Indica un camino de salida preferencial para una determinada ruta destinada a una red externa al AS;
- *ATOMIC_AGGREGATE* – Es bien conocido y discrecional. Indica si caminos más específicos se agregaron en menos específicos.
- *AGGREGATOR* - Es opcional y transitivo. Indica el ASN del último router que formó una ruta agregada, seguido por su propio ASN y dirección IP.

Atributos BGP

- Los atributos se consideran si se conoce el camino, si hay conectividad, si es accesible o si está disponible el *next hop*.
- La forma de selección puede variar dependiendo de la implementación.
- *LOCAL_PREFERENCE* es un atributo extremadamente poderoso para influenciar el tráfico de salida.
- El valor de *LOCAL_PREFERENCE* es válido para todo el AS.



Selección del camino en BGP (CISCO).

Al decidir la mejor ruta, los atributos se consideran si se conoce el camino, si hay conectividad, si es accesible o si está disponible el *next hop*. Pero la forma de selección puede variar dependiendo de la implementación.

Un atributo que vale la pena destacar es el atributo *LOCAL_PREFERENCE*. Éste es un atributo extremadamente poderoso para influenciar el tráfico de salida. Su valor es válido para todo el AS, siendo re-transmitido solamente en las sesiones iBGP.

BGP Multiprotocolo

- *Multiprotocol BGP (MP-BGP)* – Extensión de BGP para soportar múltiples protocolos de red o familias de direcciones.
 - El soporte para MP-BGP es fundamental para realizar el enrutamiento externo IPv6, ya que no existe una versión específica de BGP para esta tarea.
- Se introdujeron dos nuevos atributos:
 - *Multiprotocol Reachable NLRI (MP_REACH_NLRI)* – Transporta el conjunto de destinos alcanzables junto con la información del *next-hop*;
 - *Multiprotocol Unreachable NLRI (MP_UNREACH_NLRI)* – Transporta el conjunto de destinos inalcanzables;
 - Estos atributos son opcionales y no transitivos.

Se definieron extensiones para BGP-4 con la intención de habilitarlo para que transporte información de enrutamiento de múltiples protocolo de Capa de Red (ex., IPv6, IPX, L3VPN, etc.). Para realizar el enrutamiento externo IPv6 es fundamental el soporte para MP-BGP, ya que no existe una versión específica de BGP para tratar esta tarea.

Para que BGP pueda trabajar con la información de enrutamiento de diversos protocolos se introdujeron dos nuevos atributos:

- *Multiprotocol Reachable NLRI (MP_REACH_NLRI)*: Transporta el conjunto de destinos alcanzables junto con la información del *next-hop*;
- *Multiprotocol Unreachable NLRI (MP_UNREACH_NLRI)*: Transporta el conjunto de destinos inalcanzables.

Estos atributos son opcionales y no transitivos; en caso que un router BGP no soporte MBGP, debe ignorar esta información y no transferirla a sus vecinos.

Más información:

- RFC 2545 - *Use of BGP-4 Multiprotocol Extensions for IPv6 Inter-Domain Routing*
- RFC 4760 - *Multiprotocol Extensions for BGP-4*

BGP Multiprotocolo

- MP_REACH_NLRI
 - *Address Family Identifier* (2 bytes)
 - *Subsequent Address Family Identifier* (1 byte)
 - *Length of Next Hop Network Address* (1 byte)
 - *Network Address of Next Hop* (variable)
 - *Reserved* (1 byte)
 - *Network Layer Reachability Information* (variable)
- MP_UNREACH_NLRI
 - *Address Family Identifier* (2 bytes)
 - *Subsequent Address Family Identifier* (1 byte)
 - *Withdrawn Routes* (variable)

Estos atributos llevan la siguiente información:

MP_REACH_NLRI

- *Address Family Identifier* (2 bytes) – Identifica el protocolo de red a ser soportado;
- *Subsequent Address Family Identifier* (1 byte) – Identifica el protocolo de red a ser soportado;
- *Length of Next Hop Network Address* (1 byte) – Valor que expresa la longitud del campo *Network Address of Next Hop*, medida en bytes;
- *Network Address of Next Hop* (variable) – Contiene la dirección del siguiente salto;
- *Reserved* (1 byte) – Reservado;
- *Network Layer Reachability Information* (variable) – Lista la información de las rutas accesibles.

MP_UNREACH_NLRI

- *Address Family Identifier* (2 bytes) – Identifica el protocolo de red a ser soportado;
- *Subsequent Address Family Identifier* (1 byte) – Identifica el protocolo de red a ser soportado;
- *Withdrawn Routes* (variable) – Lista la información de las rutas inaccesibles.

Códigos más habituales para AFI y Sub-AFI

Código AFI	Código Sub-AFI	Significado
1	1	IPv4 Unicast
1	2	IPv4 Multicast
1	3	IPv4 based VPN
2	1	IPv6 Unicast
2	2	IPv6 Unicast e IPv6 Multicast RPF
2	3	Multicast RPF
2	4	IPv6 Label
2	128	IPv6 VPN
....		...

Tabla BGP

- La información sobre las rutas de Internet se encuentra en la tabla BGP.
- En los routers de borde esta información se replica hacia la RIB y la FIB, IPv4 e IPv6.
 - Tabla Global IPv4 → ~300.000 entradas
 - Tabla Global IPv6 → ~2.500 entradas
- La duplicidad de esta información implica más espacio, más memoria y más procesamiento.
 - Agregación de rutas
 - Evitar el anuncio innecesario de rutas
 - Limitar el número de rutas recibidas de otros AS
 - Importante en IPv4
 - Fundamental en IPv6

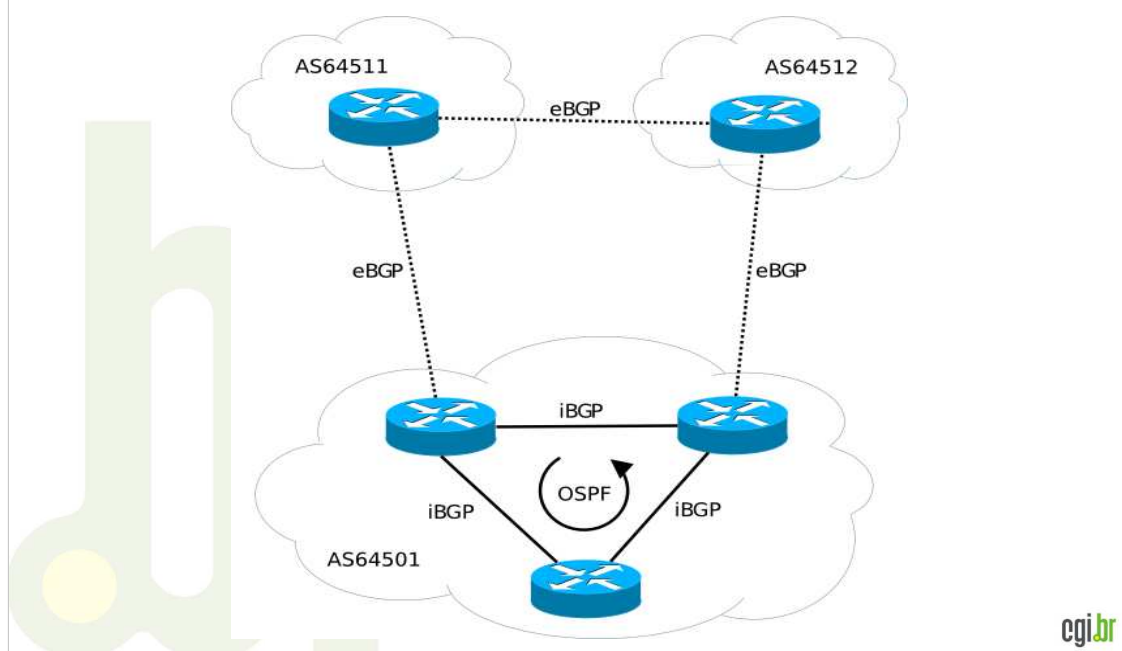
La información sobre las rutas de Internet se encuentra en la tabla BGP. En los routers de borde, los cuales se ocupan de la comunicación entre ASs, esta información se replica hacia la RIB y la FIB, IPv4 e IPv6.

La tabla global IPv4 hoy tiene aproximadamente 300.000 entradas. La tabla IPv6 tiene aproximadamente 2.500 entradas. La duplicidad de esta información en las arquitecturas distribuidas implica la necesidad de contar con más espacio para almacenamiento, más memoria y más capacidad de procesamiento, tanto en el módulo central como en las placas de las interfaces.

Esto también implica otro aspecto importante, la necesidad de establecer un plan de direccionamiento jerárquico para minimizar la tabla de rutas y optimizar el enrutamiento, evitando el anuncio de rutas innecesarias y desagregadas.

Los AS también pueden controlar los anuncios que reciben de sus vecinos BGP. Por ejemplo, es posible limitar el tamaño de los prefijos recibidos entre /20 y /24 IPv4, y entre /32 y /48 IPv6. Sin embargo, recuerde que se pueden anunciar hasta 31 prefijos IPv4 (considerando anuncios entre un /20 y un /24) y 131.071 prefijos IPv6 (considerando anuncios entre un /32 y un /48), de este modo hay quienes también controlan la cantidad de prefijos que reciben de sus vecinos BGP a través de comandos como `maximum-prefix` (Cisco) y `maximum-prefixes` (Juniper). Prestar atención a este tema es muy importante en las redes IPv4, pero en las redes IPv6 es fundamental.

Establecimiento de sesiones BGP



En este diagrama podemos analizar las opciones de configuración presentadas hasta el momento.

IPv6.br

La nueva generación del
Protocolo de Internet

246

cgi.br

Buenas Prácticas de BGP

Módulo 9

En este módulo veremos algunos conceptos básicos sobre cómo se establecen las sesiones BGP; las ventajas de utilizar interfaces de *loopback* en sesiones iBGP y eBGP; algunos aspectos de seguridad importantes que deben ser observados en la comunicación entre sistemas autónomos; formas de garantizar la redundancia y el balanceo de tráfico; además detallaremos una serie de comandos útiles para verificar el estado de las sesiones BGP.

Todos estos temas serán abordados utilizando como base las plataformas Cisco, Quagga y Juniper, y utilizando como ejemplo la configuración IPv4 implementada en nuestro laboratorio. De este modo, facilitaremos la comparación entre la configuración y los comandos utilizados con IPv6 e IPv4.

Establecimiento de sesiones BGP

- Una sesión BGP se establece entre dos routers en base a una conexión TCP.
 - puerto TCP 179;
 - conexión IPv4 o IPv6.
- Interfaz *loopback*
 - interfaz lógica;
 - no “caen”.

Una sesión BGP se establece entre dos routers en base a una conexión TCP, utilizando como estándar el puerto TCP 179, precisándose para ello conectividad IP, ya sea IPv4 o IPv6.

Una forma de establecer esa comunicación es a través de interfaces *loopback*. Éstas son interfaces lógicas, como la “null0”, es decir, “no caen” a no ser que se apague el router o que la interfaz se desconfigure.

Establecimiento de sesiones BGP

iBGP entre *loopbacks*

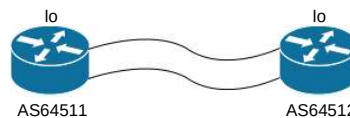
- Es fundamental establecer sesiones iBGP utilizando la interfaz *loopback*.
 - a través de la IP de la interfaz real:
 - si el *link* se interrumpe, la sesión también se interrumpirá.
 - por medio de la IP de la interfaz *loopback*:
 - más estabilidad;
 - las IP de las interfaces *loopback* serán aprendidas a través del protocolo IGP.
 - si el *link* se interrumpe, la sesión puede ser establecida por otro camino.

Por sus características, es fundamental que las sesiones iBGP se establezcan utilizando la interfaz *loopback*. En caso que la sesión sea establecida a través de la IP de la interfaz física, real, si el *enlace* se interrumpe la sesión también se interrumpirá. Si se establece a través de la interfaz *loopback*, la sesión podrá ser restablecida por otro camino aprendido a través de los protocolos IGP.

El uso de interfaces *loopback* en el establecimiento de las sesiones iBGP proporcionan mayor estabilidad a los sistemas autónomos.

Establecimiento de sesiones BGP

eBGP entre *loopbacks*



- Balanceo
- Por ejemplo:
 - Hay dos routers y cada router representa un AS:
 - Ambos están conectados mediante dos *links*;
 - Utilizando la IP de las interfaces reales:
 - Se necesitarán dos sesiones BGP;
 - Eventualmente con políticas diferentes.
 - Utilizando la IP de las interfaces *loopback*
 - Se establece una única sesión BGP;
 - Se crea una ruta estática para la IP de la interfaz *loopback* del vecino a través de cada *link*.

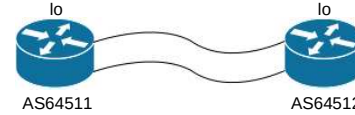
También se recomienda utilizar interfaces *loopback* para establecer sesiones eBGP. Uno de los propósitos de establecer este tipo de conexión es garantizar el balanceo.

Consideremos el siguiente ejemplo:

- Hay dos routers – cada uno de los cuales representa un AS – conectados a través de dos enlaces;
- Si en esta comunicación se utiliza la IP de las interfaces reales será necesario establecer dos sesiones BGP para cada sesión, eventualmente habrá una política diferente. Eso puede ocasionar complicaciones innecesarias.
- Utilizar las interfaces *loopback* simplifica este proceso. En ese caso, solo se necesitaría una sesión BGP y la creación de rutas estáticas apuntando a la IP de la *loopback* del vecino a través de cada enlace.

Establecimiento de sesiones BGP

eBGP entre *loopbacks*



- ¿Esa ruta estática debe ser vía interfaz o vía IP?
- Si es vía una interfaz serial se puede apuntar la ruta hacia la interfaz;
- Si es vía una interfaz Ethernet se debe apuntar a la IP.
- En un *link* serial, el tamaño de red IPv4 normalmente utilizado es /30.
 - Un /30 tiene 4 IPs: red, *broadcast*, y los dos extremos;
 - En los *links* seriales se pueden usar /31.
- ¿Cuál es el equivalente a un /31 en IPv6?
- En IPv6 se puede trabajar con redes /64 en *links* seriales.
- Una buena opción consiste en trabajar con /112.

Para establecer la sesión eBGP a través de la *loopback*, ¿la ruta estática se debe crear vía interfaz o vía IP?

Si es vía una interfaz serial se puede apuntar la ruta hacia la interfaz. Las interfaces seriales son "tubos"; la información que viaja a través de las mismas llega directamente al otro extremo, por lo tanto se puede apuntar la ruta a la interfaz.

Si se trata de un medio compartido, como por ejemplo una interfaz Ethernet, se debe apuntar a la IP.

Otro punto importante es el tamaño de red utilizado en estos tipos de enlaces. En un *enlace* serial, normalmente se utilizan prefijos de red IPv4 /30. De ese modo hay cuatro direcciones IP posibles: la de la red; la de *broadcast*; y las dos que van a identificar las interfaces. Sin embargo, en un *enlace* serial no son necesarias las direcciones de red ni de *broadcast*, y por eso hay un enfoque que utiliza prefijos de red /31 en este tipo de *enlace*. Este método también permite ahorrar una gran cantidad de direcciones IPv4, principalmente en el caso de los operadores que trabajan con miles de *enlaces* punto a punto. Por este motivo esto solo se recomienda para los *enlaces* seriales, no para los Ethernet.

En redes IPv6, el equivalente a un IPv4 /31 sería un /127. La RFC 3627 no recomienda la utilización de /127 debido a posibles problemas con la dirección *anycast Subnet-Router*, no obstante lo cual existe un *borrador* que cuestiona este argumento.

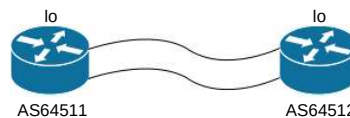
Existen diferentes posibilidades para trabajar en IPv6. En los enlaces punto a punto se puede usar un prefijo /64 o /126, pero una opción interesante consiste en utilizar un /112, de modo de trabajar apenas con el último grupo de bytes de la dirección.

Más información:

- RFC 3021 - *Using 31-Bit Prefixes on IPv4 Point-to-Point Links*
- RFC 3627 - *Use of /127 Prefix Length Between Routers Considered Harmful*
- draft-kohno-ipv6-prefixlen-p2p-00.txt - *Use of /127 IPv6 Prefix Length on P2P Links Not Considered Harmful*

Establecimiento de sesiones BGP

eBGP entre *loopbacks*



- Normalmente en la interfaz *loopback* se utiliza un prefijo /32 IPv4 o /128 IPv6.
- La IP de la *loopback* es responsabilidad del propio AS.
 - No se debe utilizar una IP privada.
- La IP del *enlace* de tránsito es responsabilidad del Proveedor de Tránsito.
 - ¿Esa IP debe o puede ser ruteable?
 - Si es una conexión IP interna con el operador, puede ser una IP válida del operador y no ruteable, por ejemplo una conexión MPLS;
 - Si es un servicio de Internet, la IP DEBE ser ruteable.

En relación al direccionamiento de las interfaces *loopback*, normalmente se utilizan prefijos IPv4 /32 y IPv6 /128. Esto desmitifica la idea de que solo es posible utilizar prefijos /64. El uso de /64 solo se obligatorio cuando se utiliza el Protocolo de Descubrimiento de Vecinos (Neighbor Discovery Protocol).

La dirección IP utilizada en la *loopback* debe formar parte del bloque del propio AS y debe ser una dirección válida, no pudiéndose utilizar IPs privadas. Las direcciones IP privadas son para uso interno en su red y la comunicación entre sistemas autónomos es una conexión a Internet, lo que exige IPs válidas.

La dirección de la interfaz real, conectada al *enlace* de tránsito, debe ser del bloque del operador que proporciona el servicio, del *Upstream Provider*. Además, esa dirección IP debe ser ruteable. Este es un tema bastante polémico, ya que hay quienes sostienen que esa dirección no debe ser ruteable por motivos de seguridad. Si hay una conexión IP interna con el operador, por ejemplo una conexión MPLS, es interesante utilizar una IP válida del operador que no sea ruteable. Sin embargo, si se trata de un servicio de Internet, el operador tiene la obligación de proporcionar una dirección ruteable ya que tener conectividad es un punto fundamental para cualquier servicio de Internet.

Ejemplo 1 – El tráfico generado por un router sale con la IP de la interfaz de salida como IP de origen. Por lo tanto, si su router se conecta a un servidor NTP, ya sea IPv4 o IPv6, la IP de origen de ese paquete será la IP de la interfaz por la cual ella sabe llegar al destino. Si es una una IP no ruteable, el paquete va a llegar al servidor pero el servidor no sabrá cómo devolver una respuesta.

Ejemplo 2 - Hay quienes piensan que el operador no debe enrutar esa IP por razones de seguridad. Si las IP internas son ruteables no hay ninguna protección, pues si hay tan solo una IP ruteable ésta será conocida por el mundo y habrá otro camino para llegar hasta ella, por ejemplo, entrando en otro elemento del AS que conozca la IP del router.

Establecimiento de sesiones BGP

eBGP entre *loopbacks*



- Seguridad
- La utilización de interfaces *loopbacks* en sesiones eBGP no es necesaria solo para asegurar el balanceo.
- Establecer sesiones eBGP utilizando la IP de la interfaz facilita mucho los ataques contra la infraestructura.
- Es recomendable establecer eBGP entre *loopbacks* aunque haya un único *enlace*.

Hay quienes sostienen que el uso de interfaz *loopback* solo es realmente necesaria para garantizar el balanceo de tráfico. Esa práctica también ayuda en relación con los temas de seguridad.

Establecer sesiones eBGP utilizando la IP de la interfaz puede facilitar los ataques contra la infraestructura de un AS. Por eso es recomendable trabajar sesiones eBGP entre *loopbacks* aunque exista un único *enlace*.

Establecimiento de sesiones BGP

eBGP entre *loopbacks*



- Seguridad
- Por ejemplo:
 - Para establecer una sesión eBGP sobre TCP se requieren 4 informaciones básicas:
 - 2 IPs y 2 puertos TCP (179 y >1024).
 - Si la sesión eBGP se establece utilizando la IP de la interfaz:
 - normalmente se identifica una de las dos utilizando *traceroute*;
 - descubriendo la primera se descubre la segunda, ya que habitualmente se utilizan /30.
 - la tercera información es un puerto estándar, el 179.
 - Es decir, 3 de las 4 variables se pueden descubrir de forma relativamente fácil.

Esto se puede ejemplificar con el siguiente escenario:

- Una sesión eBGP entre dos routers se establece a través de una conexión TCP;
- Para ello se requieren cuatro variables básicas: dos direcciones IP y dos puertos TCP, uno estándar (el 179), y al igual que en el caso de una aplicación HTTP, el router que inicia la sesión BGP va a salir por un puerto alto, superior al 1024, para cerrar con el 179;
- Si la sesión eBGP se establece utilizando la IP de la interfaz, es posible identificar esa IP utilizando comandos como el *traceroute*. Como normalmente se utilizan prefijos /30 IPv4, una vez que se descubre una IP descubrir la segunda es mucho más fácil. De este modo, dos de las cuatro variables son fáciles de descubrir;
- La tercera información es un puerto estándar, el 179.
- Es decir que de cuatro variables tres pueden ser descubiertas de manera relativamente fácil y descubrir la única variable que falta, el puerto alto, tampoco presenta mayor dificultad.

Establecimiento de sesiones BGP

eBGP entre *loopbacks*



- Seguridad
 - Una de las formas de derribar un AS o un destino consiste en derribar el AS que le proporciona conectividad.
 - Establecer una sesión eBGP utilizando *loopbacks*:
 - Las IP son de las redes internas, no tienen relación entre sí;
 - dificulta el descubrimiento mediante traceroute.

Una forma de “derribar” un AS o un destino consiste en “derribar” el AS que le provee conectividad, y eso se puede hacer interrumpiendo las sesiones eBGP.

Establecer la sesión eBGP utilizando la sesión *loopback* plantea algunos puntos relacionados con la seguridad. Las IP de la *loopback* son IPs de red interna por lo que no pueden ser descubiertas fácilmente con *tracerooutes*, y el hecho de que las dos IPs sean totalmente diferentes, sin ninguna relación entre sí, lo hace aun más difícil.

Establecimiento de sesiones BGP

- También se recomienda trabajar con una *loopback* por función y no una *loopback* por router:
 - se puede configurar una *loopback* para el Router ID, una para iBGP y una para eBGP;
 - facilita la migración de servicios;
 - aporta flexibilidad, pero consume más direcciones IP.



Otro aspecto que se debe destacar en relación con la utilización de la interfaz *loopback* es que es preferible trabajar con una *loopback* por función y no con una única *loopback* por router. Por ejemplo, se puede configurar una *loopback* para el Router ID, una para iBGP y una para eBGP.

Por ejemplo:

- Se ha establecido una sesión eBGP y es necesario migrar esa sesión para otro router;
- Si además de ser la *loopback* usada para la sesión eBGP también se utiliza para otras funciones. En este caso la migración será más complicada;
- Si se utilizara una *loopback* por función, la migración se podría realizar sin interferir con las demás funciones. Se puede migrar iBGP, el Router ID, modificar la sesión eBGP de un router a otro, sin tener que avisar al operador ni tener que migrar otros servicios internos.

A pesar de consumir más direcciones IP, esta práctica proporciona mayor flexibilidad. Sin embargo, como normalmente se utilizan prefijos /32 o /128, este problema no es tan grave.

Uso de MD5

- Una importante técnica de protección es la utilización de MD5 para la autenticación de las sesiones BGP.
- Garantiza que solamente routers confiable establezcan sesiones BGP con el AS.
- El algoritmo MD5 crea un *checksum* codificado que se incluye en el paquete transmitido.
- El router que recibe el paquete utiliza una clave de autenticación para verificar el *checksum*.
 - neighbor "ip-address o peer-group-name" password "contraseña" (Cisco)
 - authentication-key "contraseña" (Juniper)

257

Una importante técnica de protección es la utilización de MD5 para la autenticación de las sesiones BGP. De esta forma se garantiza que solamente routers confiable establezcan sesiones BGP con el AS.

El algoritmo MD5 crea un *checksum* codificado que se incluye en el paquete transmitido y el router que recibe el paquete utiliza una clave de autenticación para verificar o *hash*.

Utilice los siguientes comandos para habilitar esa función:

neighbor "ip-address ou peer-group-name" password "contraseña" (Cisco)
authentication-key "contraseña" (Juniper)

Más información:

- RFC 1321 - *The MD5 Message-Digest Algorithm*
- RFC 2385 - *Protection of BGP Sessions via the TCP MD5 Signature Option*

TTL-Security Check

- Trabajar con TTL o *Hop-Limit* igual a 1 contribuye a la seguridad
- Permite que solamente se reciban mensajes eBGP de quienes están directamente conectados;
- Pero esto es fácilmente burlado.

- La RFC5082 recomienda utilizar TTL igual a 255.

- Por ejemplo:

```
router-R13(config-router)# neighbor 2001:DB8:200:FFFF::255  
ttl-security hops 1
```

- Define el valor mínimo esperado para el *Hop-Limit* de entrada como al menos 254 (255 - 1).
- El router aceptará la sesión a partir de 2001:DB8:200:FFFF::255 si está a 1 salto de distancia.

258

Trabajar con TTL o *Hop-Limit* igual a 1 contribuye a la seguridad de las sesiones BGP, ya que que solamente se reciban mensajes eBGP de quienes están directamente conectados. Pero esto es burlado fácilmente, alcanza con utilizar comandos como **traceroute** para identificar cuántos saltos son necesarios para llegar al router de destino y generar un paquete con el valor del TTL necesario para alcanzarlo.

La RFC 5082 recomienda trabajar con TTL igual a 255 en lugar de 1. De este modo, en una sesión eBGP directamente conectada utilizando la IP de la interfaz se puede garantizar que el vecino BGP está como máximo a un salto de distancia a través de la lectura del valor del TTL, que ha sido decrementado en cada salto.

Para utilizar esta funcionalidad es necesario configurar los dos vecinos que participan en la sesión eBGP. Quien envíe el mensaje debe armar el paquete con TTL igual a 255 y quien lo recibe debe habilitar la verificación de este campo. En los routers Cisco la verificación se puede realizar de la siguiente manera:

```
router-R13(config-router)#neighbor 2001:DB8:200:FFFF::255  
ttl-security hops 1
```

De este modo se define el valor mínimo esperado para el *Hop-Limit* de entrada como al menos 254 (255 - 1). Así el router aceptará la sesión a partir de 2001:DB8:200:FFFF::255 si está a 1 salto de distancia.

Con un vecino BGP IPv4 esta línea sería:

```
router-R13(config-router)#neighbor 10.2.255.255 ttl-security  
hops 1
```

Más información:

- RFC 5082 - The Generalized TTL Security Mechanism (GTSM)
- http://www.cisco.com/en/US/docs/ios/12_3t/12_3t7/feature/guide/gt_btsh.html
- <http://www.juniper.net/us/en/community/junos/script-automation/library/configuration/ttl-security/>

TTL-Security Check

- Este es el tercer mecanismo de protección de eBGP presentado hasta el momento:
 - 1º – establecer la sesión entre *loopbacks*;
 - 2º – Usar MD5;
 - 3º – Usar *TTL-Security Check*.
- El *TTL-Security Check* se utiliza poco, pero es extremadamente útil.
- Solo enviar el paquete con TTL 255 no es suficiente. También es necesario configurar el vecino.
 - la sesión eBGP puede ser establecida por un *enlace* que no sea el correcto;
 - dificulta la detección del origen de los problemas.
- Las sesiones entre *loopbacks* usan *ttl-security hops* 2.²⁵⁹

Este es el tercer mecanismo de protección de eBGP presentado hasta el momento:

- 1º – establecer la sesión entre *loopbacks*;
- 2º – Usar MD5;
- 3º – Usar *TTL-Security Check*.

El *TTL-Security Check* se utiliza poco, pero es extremadamente útil. Sin embargo, solo enviar el paquete con TTL 255 no es suficiente: también es necesario configurar el vecino. Si eso no ocurre, la sesión eBGP se podrá establecer por un enlace que no sea el correcto, lo que dificultará la detección del origen de los problemas.

Otro punto importante es que en las sesiones entre *loopbacks* se debe usar *ttl-security hops* 2.

Deshabilitar el descubrimiento de vecinos

- Hay routers que traen el anuncio de mensajes RA habilitado por defecto.
- Si en la interfaz del router se utiliza una dirección /64 va a haber descubrimiento de vecinos, incluso entre routers.
 - Con esto el router puede anunciar que es el *gateway* estándar;
 - Puede generar *looping*
- No hay problemas en los *enlaces* para estaciones de trabajo.
- En los *enlaces* entre routers se debe deshabilitar el envío de RA.
 - `ipv6 nd ra suppress` (Cisco)
 - `ipv6 nd suppres-ra` (Cisco / Quagga / Juniper)

260

Un punto importante en la configuración de routers IPv6 es el funcionamiento del protocolo de Descubrimiento de Vecinos (Neighbor Discovery). Hay routers que traen el anuncio de mensajes RA (*Router Advertisements*) habilitado por defecto. Si en la interfaz del router se utiliza una dirección /64 va a haber descubrimiento de vecinos, incluso entre routers. De este modo el router puede anunciar que es el *gateway* estándar para los demás routers de la red, pudiendo generar *loopings*.

Esta función solo debe ser habilitada en interfaces que estén conectadas a estaciones de trabajo o en algunos casos a servidores. En los enlaces entre routers se debe deshabilitar el envío de RA.

Esta función se puede deshabilitar utilizando los siguientes comandos:

`ipv6 nd ra suppress` (Cisco)

`ipv6 nd suppres-ra` (Cisco / Quagga / Juniper)

Verificación de las configuraciones

- Verificar los protocolos configurados:
 - `show ip protocols` (Cisco)
 - `show ipv6 protocols` (Cisco)
 - En Quagga existe un *daemon* específico para cada protocolo de enrutamiento, los cuales son tratados como procesos separados.
- Verificar el estado y las direcciones de las interfaces:
 - `show ip interface brief` (Cisco)
 - `show ipv6 interface brief` (Cisco)
 - `show interface terse` (Juniper v4 y v6)
 - Observe que, en caso de trabajar con sub-interfaces, la dirección IPv6 *link-local* será la misma. Son interfaces lógicas distintas, pero la dirección está compuesta por la MAC de la física.

261

Conocer las funciones y configuraciones habilitadas en los routers es muy importante, particularmente cuando se adquiere un equipo nuevo. Estos son algunos comandos que pueden facilitar esta tarea:

Para verificar los protocolos configurados podemos utilizar:

`show ip protocols` (Cisco)

`show ipv6 protocols` (Cisco)

En Quagga existe un *daemon* específico para cada protocolo de enrutamiento, los cuales son tratados como procesos separados.

Verificación del estado y las direcciones de las interfaces:

`show ip interface brief` (Cisco)

`show ipv6 interface brief` (Cisco)

`show interface terse` (Juniper v4 y v6)

Observe que, en caso de trabajar con sub-interfaces, la dirección IPv6 *link-local* será la misma. Son interfaces lógicas distintas, pero la dirección está compuesta por la MAC de la física.

Estos comandos son útiles porque muestran de forma resumida lo que está configurado en el equipo, por ejemplo qué interfaces ya tienen IPv6 habilitado, etc.

Verificación de las configuraciones de BGP y de iBGP

- Visualización de la configuración actual de BGP (Cisco):

```
router-R13#show running-config | begin bgp
router bgp 64501
no bgp default ipv4-unicast
bgp log-neighbor-changes
neighbor 2001:DB8:21:FFFF::254 remote-as 64501
neighbor 2001:DB8:21:FFFF::254 description R12
neighbor 2001:DB8:21:FFFF::254 update-source Loopback20
neighbor 2001:DB8:21:FFFF::254 version 4
neighbor 2001:DB8:21:FFFF::255 remote-as 64501
neighbor 2001:DB8:21:FFFF::255 description R11
neighbor 2001:DB8:21:FFFF::255 update-source Loopback20
neighbor 2001:DB8:21:FFFF::255 version 4
neighbor 2001:DB8:200:FFFF::255 remote-as 64512
neighbor 2001:DB8:200:FFFF::255 description R03
neighbor 2001:DB8:200:FFFF::255 ebgp-multihop 2
neighbor 2001:DB8:200:FFFF::255 update-source Loopback30
neighbor 2001:DB8:200:FFFF::255 version 4      262
...
```

Existen algunos comandos que pueden facilitar el análisis y la verificación de las políticas de entrada y salida de un router.

Para visualizar la configuración actual de BGP en un router Cisco o Quagga:

```
show running-config | begin bgp
```

En nuestro ejemplo, la primera línea de la configuración de BGP indica el ASN del propio AS:

- `router bgp 64501`

Por defecto, los routers Cisco y Quagga solo conocen una familia de direcciones, la familia *ipv4-unicast*. Para utilizar otras familias de direcciones se utiliza el concepto de *address-family*, y para habilitarlo se utiliza el comando:

- `no bgp default ipv4-unicast`

Para habilitar IPv6 en los routers Cisco y Quagga se recomienda marcar una ventana de mantenimiento para de este modo poder interrumpir el tráfico, aplicando el comando

- `no router bgp 64501`

y rehacer toda la configuración con *address-family*.

Incluso utilizando *address-family*, al inicio de la configuración siempre se presentan los datos generales que no dependen de la familia. Por ejemplo:

- `neighbor 2001:DB8:200:FFFF::255 remote-as 64512` – indica el IP y el ASN del vecino. Si ASN indicado es el del propio AS es porque se trata de una sesión iBGP;
- `neighbor 2001:DB8:200:FFFF::255 description R03` – presenta un nombre de identificación;

- `neighbor 2001:DB8:200:FFFF::255 ebgp-multihop 2` – especifica el número de saltos hasta llegar al vecino. Una diferencia importante entre iBGP y eBGP es que cuando el router genera un mensaje eBGP el paquete IP que lleva ese mensaje es enviado con el valor del TTL, si es IPv4, o del *Hop_Limit*, si es IPv6, igual a 1, y por lo tanto solo puede realizar un salto. Si un paquete tiene que ser enrutado a una interfaz *loopback* y el mensaje sale con el TTL igual a 1, al abrirlo el router de destino decrementará el valor del TTL y no podrá realizar el enrutamiento interno a la *loopback*. Por lo tanto, para levantar la sesión eBGP entre *loopbacks* se debe especificar cuál es el número de saltos;

- `neighbor 2001:DB8:200:FFFF::255 update-source Loopback30` - configura el router para que la IP de origen de los paquetes sea la de la *loopback*, porque al enviar un mensaje el router adopta como dirección IP de origen por defecto la IP de la interfaz por donde será enviada.

- `neighbor 2001:DB8:200:FFFF::255 version 4` – indica la versión del protocolo BGP utilizado. Esta información agiliza el establecimiento de la sesión BGP, ya que en el primer mensaje intercambiado entre los vecinos se transmiten algunos datos, entre ellos la negociación de la versión. Si la versión utilizada se informa desde el inicio no es necesario realizar esta negociación.

En cuanto a la configuración de iBGP, las únicas diferencias son que el ASN es el del propio AS y que, a diferencia de eBGP, no es necesario cambiar el TTL. Por defecto, en iBGP el TTL no se modifica, en el supuesto de que la sesión puede hacer un largo camino, saliendo con TTL igual a 255 u otro valor intermedio dependiendo de la implementación.

A continuación presentamos un ejemplo de las configuraciones de una sesión BGP IPv4:

```
router-R13#show running-config | begin bgp
router bgp 64501
  no bgp default ipv4-unicast
  bgp log-neighbor-changes
  neighbor 10.2.255.255 remote-as 64512
  neighbor 10.2.255.255 description R03
  neighbor 10.2.255.255 ebgp-multihop 2
  neighbor 10.2.255.255 update-source Loopback30
  neighbor 10.2.255.255 version 4
  neighbor 172.21.15.254 remote-as 64501
  neighbor 172.21.15.254 description R12
  neighbor 172.21.15.254 update-source Loopback20
  neighbor 172.21.15.254 version 4
  neighbor 172.21.15.255 remote-as 64501
  neighbor 172.21.15.255 description R11
  neighbor 172.21.15.255 update-source Loopback20
  neighbor 172.21.15.255 version 4
  ...
```


Configuración del *address-family*

- En los routers Cisco y Quagga, para utilizar IPv6 es necesario especificar la familia de direcciones con la cual se está trabajando.
- Aplicar las configuraciones específicas de cada familia para cada vecino

```
router-cisco# show running-config | begin address-family ipv6
address-family ipv6
neighbor 2001:DB8:21:FFFF::254 activate
neighbor 2001:DB8:21:FFFF::254 next-hop-self
neighbor 2001:DB8:21:FFFF::254 soft-reconfiguration inbound
neighbor 2001:DB8:21:FFFF::255 activate
neighbor 2001:DB8:21:FFFF::255 next-hop-self
neighbor 2001:DB8:21:FFFF::255 soft-reconfiguration inbound
neighbor 2001:DB8:200:FFFF::255 activate
neighbor 2001:DB8:200:FFFF::255 soft-reconfiguration inbound
neighbor 2001:DB8:200:FFFF::255 route-map BGPIn-IPv6-AS64512 in
neighbor 2001:DB8:200:FFFF::255 route-map BGPout-IPv6-AS64512 out
network 2001:DB8:21::/48
network 2001:DB8:21:8000::/49
exit-address-family                                     264
```

En los routers Cisco y Quagga, para utilizar IPv6 es necesario especificar la familia de direcciones con la cual se está trabajando. A diferencia de los routers Juniper, la configuración de BGP se presenta dividida en configuraciones generales y configuraciones específicas de cada familia para cada vecino.

Para analizar las configuraciones del *address-family* IPv6 se utiliza:

```
show running-config | begin address-family ipv6
```

- `address-family ipv6` – indica a cuál familia pertenece las configuración;
- `neighbor 2001:DB8:200:FFFF::255 activate` – activa la sesión, necesario cuando se trabaja con *address-family*. Una buena práctica al configurar una sesión iBGP o eBGP consiste en levantarla en *shutdown*. Esto evita que la sesión se establezca sin que las políticas estén configuradas, no permitiendo el envío de datos indebidos;
- `neighbor 2001:DB8:200:FFFF::255 soft-reconfiguration inbound` – indica la forma en que la tabla de enrutamiento será actualizada;
- `neighbor 2001:DB8:200:FFFF::255 prefix-list BGPout-IPv6-AS64512 out` – indica la política de salida aplicada;
- `neighbor 2001:DB8:200:FFFF::255 route-map BGPIn-IPv6-AS64512 in` – indica la política de salida aplicada.

En cuanto a la configuración de iBGP se destaca la siguiente información:

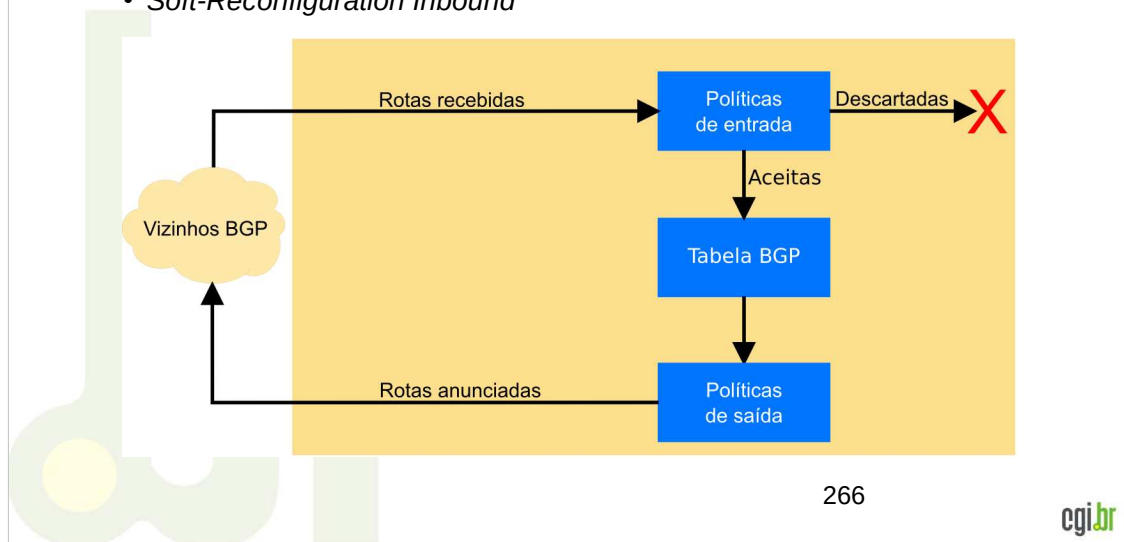
- `neighbor 2001:DB8:21:FFFF::254 next-hop-self` – indica que es el siguiente salto. Esta configuración aporta más estabilidad y facilita la operación de los AS. Un router de borde puede transmitir a los demás routers de su AS, via iBGP, todos los prefijos que aprende de sus AS vecinos. Al pasar a los demás routers se mantiene el atributo *next-hop*. Sin embargo, el *next-hop* de estos prefijos siempre será el router de borde de los vecinos, con los cuales los routers internos no poseen conectividad directa. Para solucionar este problema, el router de borde del AS transmite los anuncios informando que el *next-hop* para los AS vecinos es él mismo, a través del comando `next-hop-self`. De este modo los routers internos solo precisan saber llegar al router de borde de su AS, que es el que tiene conectividad a Internet.

A continuación presentamos un ejemplo de configuración para el *address-family* IPv4:

```
router-cisco# show running-config | begin address-family ipv4
address-family ipv4
neighbor 10.2.255.255 activate
neighbor 10.2.255.255 soft-reconfiguration inbound
neighbor 10.2.255.255 prefix-list BGPout-IPv4-AS64512 out
neighbor 10.2.255.255 route-map BGPIn-IPv4-AS64512 in
neighbor 172.21.15.254 activate
neighbor 172.21.15.254 next-hop-self
neighbor 172.21.15.254 soft-reconfiguration inbound
neighbor 172.21.15.255 activate
neighbor 172.21.15.255 next-hop-self
neighbor 172.21.15.255 soft-reconfiguration inbound
network 172.21.0.0 mask 255.255.240.0
network 172.21.8.0 mask 255.255.248.0
exit-address-family
```

Configuración del *address-family*

- *Soft-Reconfiguration Inbound*



266

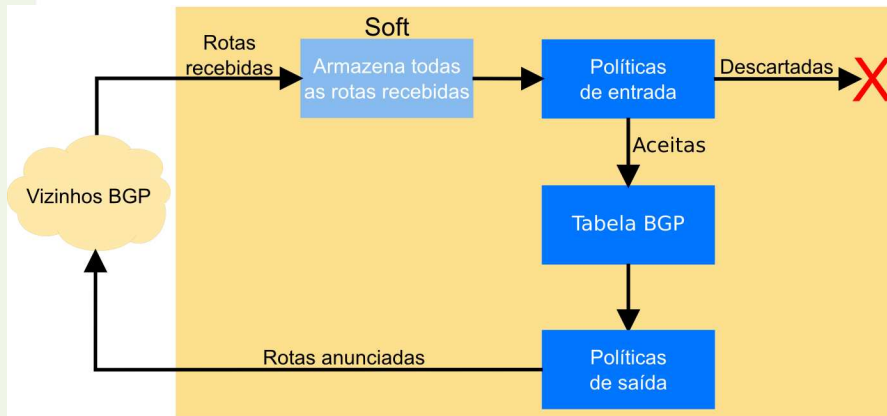
cgi.br

`soft-reconfiguration inbound` es un comando interesante. Consideremos el siguiente ejemplo:

- El router R1 levanta una sesión BGP con el router R2;
- Cuando se establece la sesión el router envía toda la información que conoce;
- Solo se enviarán nuevas informaciones cuando sea necesario agregar o quitar entradas de la tabla;
- Si en el router R2 se crea una política de entrada, el mensaje original que se intercambió al establecer la sesión será modificado;
- Si fuera necesario crear una nueva política en R2 ya no se tendrán los datos iniciales para poder aplicarlas.

Configuración del *address-family*

- *Soft-Reconfiguration Inbound*



- `router-R13# clear bgp ipv6 unicast 2001:DB8:200:FFFF::255 soft in`

Una forma de recuperar estos datos sería “derribar” la sesión para que el router envíe nuevamente todos sus prefijos. Esta práctica era útil cuando no había tantas entradas en la tabla global de enrutamiento. Pero hoy en día ya no es efectiva.

Otra opción es utilizar el comando `soft-reconfiguration inbound`. Así, antes de aplicarse las políticas se crea otra tabla de entrada por vecino, exactamente igual a la recibida. De este modo todo lo que el R1 envía será grabado en esta tabla previa, salvando los prefijos originales. Si fuera necesario modificar alguna configuración se puede utilizar, por ejemplo, el comando:

```
router-R13# clear bgp ipv6 unicast 2001:DB8:200:FFFF::255 soft in
```

Este comando hace que el router lea nuevamente la tabla previa sin interrumpir la sesión. Pero esto no funciona hoy en día porque la tabla BGP tiene alrededor de 300 mil prefijos, y con este comando se duplica la tabla BGP para cada vecino, consumiendo mucha más memoria de su módulo de enrutamiento, en la parte de control.

El siguiente sería un ejemplo de su utilización con IPv4:

```
router-R13# clear bgp ipv4 unicast 10.2.255.255 soft in
```

Configuración del *address-family*

- *Route Refresh*
 - Cuando los routers inician una sesión BGP, cada router transmite una serie de datos sobre los recursos que conoce, por ejemplo: qué *capacidades* soporta.
 - Una de ellas es el *route-refresh*.
 - Permite recuperar los datos originales de la tabla de enrutamiento sin “derribar” la sesión BGP y sin crear tablas adicionales.
 - Solicita al vecino el reenvío de la tabla de enrutamiento.
 - Para saber si el router soporta *route-refresh* use el comando:
 - `show ipv6 bgp neighbor 2001:DB8:200:FFFF::255`

268

Cuando los routers inician una sesión BGP, cada router transmite una serie de datos sobre los recursos que conoce, por ejemplo: que *capacidades* soporta. Una de ellas es el *route-refresh*.

Este recurso permite recuperar los datos originales de la tabla de enrutamiento sin “derribar” la sesión BGP y sin crear tablas adicionales, solo solicitando al vecino el reenvío de la tabla de enrutamiento.

Para saber si el router soporta *route-refresh* un ejemplo sería:

```
show ipv6 bgp neighbor 2001:DB8:200:FFFF::255
show ip bgp neighbor 10.2.255.255
```

Con este comando también se puede ver si hay soporte para otras *capacidades*, como el soporte para ASN de 32 bits (*New ASN Capability*).

Verificación de las configuraciones de BGP y de iBGP

- Visualización de la configuración actual de BGP (Juniper):

```
juniper@R11> show configuration protocols bgp
protocols {
  bgp {
    group iBGPv6 {
      type internal;
      local-address 2001:DB8:21:FFFF::255;
      export next-hop-self;
      neighbor 2001:DB8:21:FFFF::252;
      neighbor 2001:DB8:21:FFFF::254;
    }
    group eBGP-AS64511v6 {
      type external;
      neighbor 2001:db8:100:1::1 {
        import nh-BGPIn-IPv6-AS64511;
        export nh-BGPout-IPv6-AS64511;
        peer-as 64511;
      }
    }
  }
}
```

269

cgi.br

Los routers Juniper ya trabajan con el concepto de *address-family* por defecto. (inet, inet6). Para visualizar la configuración actual de BGP en un router Juniper:

show configuration protocols bgp

En el primer grupo se presentan las configuraciones de iBGP informando los siguientes datos:

- **group iBGPv6** – nombre del grupo;
- **type internal** – indica que es iBGP;
- **local-address 2001:DB8:21:FFFF::255** – dirección de la interfaz de salida;
- **export next-hop-self** – propaga a los routers internos que el siguiente salto a cualquier destino es el router de borde del propio AS;
- **neighbor 2001:DB8:21:FFFF::252** – indica la IP del vecino iBGP;
- **neighbor 2001:DB8:21:FFFF::254** – indica la IP del vecino iBGP.

El segundo grupo muestra los datos de eBGP:

- **group eBGP-AS64511** – nombre del grupo;
- **type external** – indica que es eBGP;
- **neighbor 2001:db8:100:1::1** – indica la dirección del vecino eBGP;
- **import nh-BGPIn-IPv6-AS64511** – política de entrada aplicada;
- **export nh-BGPout-IPv6-AS64511** – política de salida aplicada;
- **peer-as 64511** – ASN del vecino.

A diferencia de la configuración del router Cisco presentada anteriormente, en el ejemplo anterior utilizamos una dirección IP de la interfaz real para establecer las sesiones BGP.

A continuación presentamos un ejemplo de las configuraciones de una sesión BGP IPv4 en un router Juniper:

```
juniper@R11> show configuration protocols bgp
protocols {
  bgp {
    group iBGP {
      type internal;
      local-address 172.21.15.255;
      export next-hop-self;
      neighbor 172.21.15.252;
      neighbor 172.21.15.254;
    }
    group eBGP-AS64511 {
      type external;
      neighbor 10.1.1.1 {
        import nh-BGPin-IPv4-AS64511;
        export nh-BGPout-IPv4-AS64511;
        peer-as 64511;
      }
    }
  }
}
```

Decisiones de enrutamiento

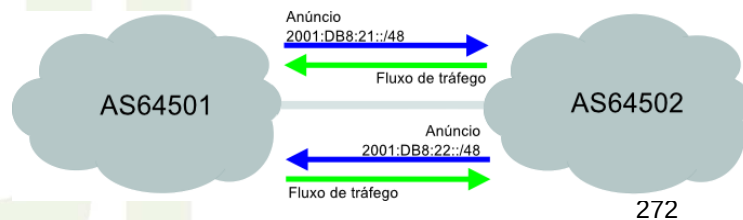
- Los routers toman decisiones de acuerdo con la información que conocen.
- Esta información es recibida y enviada a otros routers a través de los protocolos de enrutamiento interno y externo.
 - Los routers solo anuncian la mejor ruta que conocen para un determinado destino.
- Esta información se utiliza para influenciar el tráfico de entrada y de salida del AS.

Los routers toman decisiones de acuerdo con la información que conocen. Esta información es recibida y enviada a otros routers a través de los protocolos de enrutamiento interno y externo.

Al enviar su información, los routers solo anuncian la mejor ruta que conocen para un determinado destino. Esta información se utilizará para influenciar el tráfico de entrada y de salida del AS.

Influencia sobre el tráfico

- ¿Los prefijos que anuncia un AS interfieren con el tráfico de entrada o salida?
 - Los prefijos anunciados interfieren con la manera en que los demás conocen al AS.
 - tráfico de entrada.
 - Los prefijos recibidos de otras redes interfieren con el tráfico de salida.



Los prefijos anunciados interfieren con la manera en que los demás conocen al AS, es decir, interfieren con el tráfico de entrada. De la misma manera, los prefijos recibidos interfieren con el tráfico de salida.

Influencia sobre el tráfico

- ¿Qué es más fácil, influenciar el tráfico de entrada o de salida?
- Ejemplo:
 - Un AS tiene un bloque IPv4 /20;
 - Este AS puede generar para Internet anuncios de prefijos hasta un /24, el prefijo IPv4 más específico normalmente aceptado por los operadores;
 - ¿Cuántos prefijos /24 se pueden generar a partir de un /20?
 - ¿Cuántos prefijos se pueden generar entre un /20 y un /24?
 - ¿Y entre un /32 y un /48 IPv6?

273

¿Qué es más fácil en un AS, influenciar el tráfico de entrada o de salida?

Analice el siguiente ejemplo:

- Un AS tiene un bloque IPv4 /20;
 - Este AS puede generar para Internet anuncios de prefijos hasta un /24, el prefijo IPv4 más específico normalmente aceptado por los operadores;
 - Con un /20 se pueden generar 16 prefijos /24;
 - Considerando la hipótesis de anunciar todos los prefijos posibles entre el /20 y el /24, se puede generar un total de 31 prefijos;
 - Y con un bloque de direcciones IPv6, ¿cuántos prefijos se pueden generar entre un /32 y un /48?

Influencia sobre el tráfico

- Internet sabe llegar al AS hasta por 31 prefijos IPv4.
- Pero ¿cuántas entradas IPv4 de Internet conoce un AS?
- Por consiguiente, es mucho más fácil influenciar el tráfico de salida.
 - Mayor cantidad de información;
 - Las decisiones de enrutamiento se basan en los prefijos.
 - Balanceo de tráfico;
 - Contabilidad de tráfico;
 -
- La influencia del tráfico de entrada y salida está relacionada con la política de enrutamiento a ser implementada.
 - Hay dos frentes: el de entrada y el de salida, llamados AS-IN y AS-OUT.
- Igual para IPv4 e IPv6.

274

De este modo, Internet sabe llegar al AS hasta por 31 prefijos IPv4 y el AS tiene la opción de salir por aproximadamente 300.000 prefijos, que es el tamaño actual de la Tabla de Enrutamiento Global IPv4.

Por lo tanto, si pensamos que “información es poder”, podemos afirmar que es más fácil influenciar el tráfico de salida, ya que hay una mayor cantidad de prefijos para trabajar y las decisiones de enrutamiento se basan en los prefijos, actuando sobre el balanceo, contabilidad del tráfico, etc.

La influencia de los tráficos de entrada y salida está relacionada con la política de enrutamiento a ser implementada, que tiene dos frentes: el de entrada y el de salida, llamados AS-IN y AS-OUT. Esto ocurre de la misma manera para IPv4 y para IPv6.

Plan de direccionamiento

- Distribución de los servicios, servidores, etc., entre diferentes partes del bloque de direcciones IP.
- facilita la influencia del tráfico de entrada y salida de su AS;
- no es bueno concentrar todo el tráfico principal detrás del mismo prefijo /24 o /48 anunciado en Internet.
- ¿Esta mala distribución limitará la influencia del tráfico de entrada o de salida?

Un punto importante del plan de direccionamiento IP, tanto IPv4 como IPv6, es la distribución de los, servidores, etc., entre diferentes partes del bloque de direcciones IP, de modo de facilitar la influencia del tráfico de entrada y de salida del AS.

No es recomendable ubicar todos los servidores, clientes importantes que generan la mayor parte del tráfico, en el mismo /24 IPv4 o /48 IPv6. Esta mala distribución limitará la influencia del tráfico de entrada. Porque el tráfico de entrada es influenciado por los anuncios generados. Es decir que es fundamental ver cómo se ubicarán los consumidores de tráfico de entrada y de salida en el plan de direccionamiento.

Los consumidores de tráfico de entrada, es decir los usuarios de acceso a Internet, tienen que distribuirse bien entre los prefijos anunciados.

Plan de direccionamiento

- AS-OUT
 - Es lo que será anunciado en Internet;
 - Afecta el tráfico de entrada.
 - Por ejemplo:
 - El AS54501 tiene un /48 IPv6;
 - para realizar el balanceo de tráfico, haremos que la mitad de este tráfico ingrese por un *enlace* y la otra mitad por otro;
 - El /48 se divide en dos /49, anunciando el primer /49 en un *enlace* y el segundo /49 en otro .

276

AS-OUT es la política que se ocupa de lo que será anunciado en Internet. Es lo que va a interferir con el tráfico de entrada.

Por ejemplo, el AS64501 tiene un bloque /48 IPv6 y, para dividir y balancear el tráfico de entrada por dos enlaces de acceso a Internet, se lo debe influenciar para que una mitad entre por un enlace y la otra mitad entre por otro enlace. Para eso el /48 se divide en dos prefijos /49, anunciando el primer /49 en un enlace y el segundo /49 en otro. Esto se utiliza para permitir el balanceo del tráfico.

Plan de direccionamiento

- AS-IN
 - Depende de los anuncios recibidos de Internet
 - normalmente la tabla completa.
 - Afecta el tráfico de salida.
 - Se puede influenciar el tráfico de salida alterando el valor de *LOCAL_PREFERENCE* de acuerdo con determinadas condiciones.
 - *LOCAL_PREFERENCE* es el atributo con mayor fuerza para influenciar el tráfico de salida.
 - Por ejemplo: El AS54501 debe influenciar su tráfico de salida, de modo que el tráfico con destino al primer /49 del AS64513 salga preferentemente por el *enlace* con el AS64511 y el tráfico con destino al segundo /49 del AS64513 salga preferentemente por el *enlace* con el AS64512.
 - Preferentemente es una palabra clave para el BGP.²⁷⁷

AS-IN es la política que interfiere con el tráfico de salida. Ésta depende de los anuncios recibidos de Internet, normalmente de la tabla de enrutamiento completa.

Para influenciar el tráfico de salida se puede alterar el valor del atributo *LOCAL_PREFERENCE* de los anuncios recibidos, de acuerdo con determinadas condiciones. Es el atributo con mayor fuerza para influenciar el tráfico de salida.

Por ejemplo: El AS54501 necesita influenciar su tráfico de salida, de modo que el tráfico con destino al primer /49 del AS64513 salga preferentemente por el *enlace* con el AS64511 y el tráfico con destino al segundo /49 del AS64513 salga preferentemente por el *enlace* con el AS64512.

Preferentemente es una palabra clave para el BGP. Con BGP es fundamental trabajar con preferencias y no con filtros. Existen diferentes maneras de hacerlo, pero para que una configuración sea realmente efectiva, especialmente cuando ocurre una caída en un *enlace*, es importante no descartar nada, se debe dar preferencia a un camino sobre otro, pero no descartar ninguno. El descarte tiene un efecto inmediato, pero deja de ser efectivo cuando no hay redundancia.

Plan de direccionamiento

- Redundancia
- Cada /49 IPv6 es conocido por el mundo a través de un solo *enlace*
 - Si uno de esos *enlaces* se cae, el /49 anunciado por el mismo quedará inaccesible.
- Para tener redundancia también se debe anunciar el /48 en los dos *enlaces*.
- Como la preferencia es por el prefijo más específico, si los dos *enlaces* están activos Internet va a preferir los /49.
- Cuando uno de los dos *enlaces* se cae y uno de los dos /49 deja de ser anunciado, Internet aun tiene la opción del /48 anunciado en otro *enlace*, garantizando así la redundancia.
- El tráfico se debe distribuir entre ambos, colocando la mitad de los consumidores de tráfico de entrada en un /49 y la mitad en el otro.

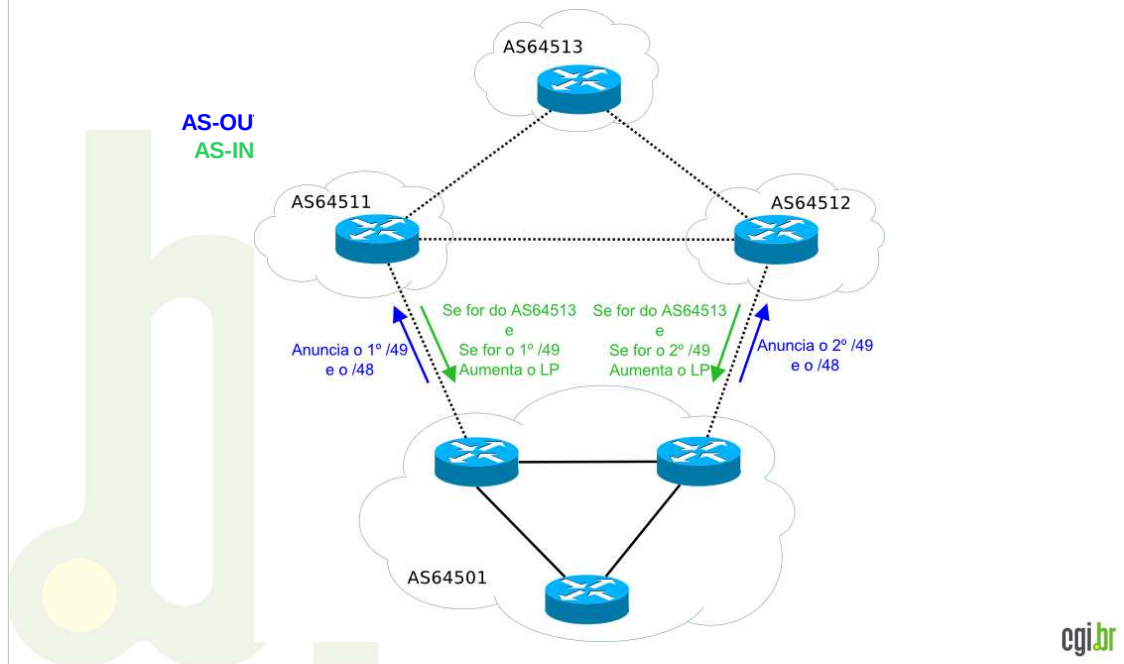
278

Para que haya redundancia se debe proceder de la siguiente manera:

- Luego de dividir el bloque /48 IPv6 en dos prefijos /49, cada uno de ellos se anuncia en un *enlace*, es decir, cada /49 es conocido por el mundo a través de un solo *enlace*, un camino;
- Si uno de esos *enlaces* se cae, el /49 anunciado por el mismo quedará inaccesible;
- Para tener redundancia también se debe anunciar el prefijo /48 en los dos *enlaces*;
- Así Internet conocerá el /48 por los dos caminos;
- Como la preferencia es por el prefijo más específico, si los dos *enlaces* están activos Internet va a preferir siempre los /49, es decir, el balanceo estará funcionando;
- Cuando uno de esos *enlaces* se cae, uno de los prefijos /49 dejará de ser anunciado. Pero este /49 está contenido en el /48, es decir que, aunque uno de los *enlaces* esté desactivado, Internet aun tendrá la opción del /48 anunciado en el otro *enlace*. Esto garantiza la redundancia.

Para realmente distribuir el tráfico entre ambos la mitad de los consumidores de tráfico de entrada se debe colocar en un /49 y la mitad en el otro. Esto forma parte de la planificación.

Plan de direccionamiento



En este diagrama podemos analizar ejemplos de las políticas de enrutamiento presentadas hasta el momento.

- AS-OUT – Utiliza los anuncios enviados para interferir con el tráfico de entrada;
- AS-IN – Utiliza los anuncios recibidos para influenciar el tráfico de salida;

Políticas de enrutamiento

- Una importante función de BGP está relacionada con la manipulación de los atributos y los tests condicionales:
 - Cisco
 - *route-map* – define las condiciones para la redistribución de rutas y permite controlar y modificar datos de las políticas de enrutamiento;
 - *prefix-list* – mecanismo de filtrado de prefijos muy poderoso. Permite trabajar con notación de prefijo, agregar descripciones y trabajar con secuencias;
 - Juniper
 - *route-filter* – utilizado para comparar rutas individualmente o en grupos.

280

Una importante función de BGP es la manipulación de los atributos y los tests condicionales: Para tratar estos aspectos tenemos las siguientes funcionalidades:

- *route-map* – define las condiciones para la redistribución de rutas y permite controlar y modificar datos de las políticas de enrutamiento;
- *prefix-list* (Cisco y Quagga) – mecanismo de filtrado de prefijos con muchos recursos. Permite trabajar con notación de prefijo, agregar descripciones y trabajar con secuencias, presentando así una ventaja con respecto al uso de la función *distribute-list*, que por ser basada en ACLs facilita el filtrado de paquetes pero es difícil de manejar;
- *route-filter* (Juniper) – utilizado para comparar rutas individualmente o en grupos.

Análisis del AS-PATH

- AS-PATH - Atributo fundamental de BGP. Consiste en el ASN de las redes por las cuales pasará el paquete hasta llegar al destino.
- Análisis del AS-PATH con expresiones regulares:

- Cisco / Quagga

```
ip as-path access-list 32 permit .*  
ip as-path access-list 69 deny .*  
ip as-path access-list 300 permit (_64513)+$
```

- Juniper

```
as-path ALL .*;  
as-path AS64513 ".*( 64513)+$";
```

281

cgi.br

BGP es un protocolo que se usa en la comunicación entre los AS. Por este motivo el AS-PATH es un atributo fundamental de BGP. Consiste en el ASN de las redes por las cuales pasará el paquete hasta llegar al destino.

Los routers Cisco, Quaga y Juniper ofrecen comandos que permiten analizar el AS-PATH usando expresiones regulares. En nuestro ejemplo, en las expresiones regulares tenemos los siguientes términos:

- El carácter “punto” significa 'qualquier elemento'
- El carácter “asterisco” significa 'zero o varias ocurrencias'
- El carácter “\$” significa 'final de línea'
- El carácter “+” significa 'una o más ocurrencias'

```
ip as-path access-list 32 permit .* (Cisco / Quagga)  
as-path ALL .*; (Juniper)
```

- Estas líneas indican que cualquier AS-PATH será permitido.

```
ip as-path access-list 69 deny .* (Cisco / Quagga)
```

- Esta línea indica que cualquier bloque será negado.

```
ip as-path access-list 300 permit (_64513)+$ (Cisco / Quagga)  
as-path AS64513 ".*( 64513)+$"; (Juniper)
```

- Estas líneas indican que todos los prefijos originados en el AS 64513 serán permitidos. Se permite una o más ocurrencias para garantizar *AS PATH prepends*, es decir repeticiones de un mismo ASN en secuencia a lo largo del AS-PATH (en este caso el ASN 64513).

Establecimiento de filtros

- Política de entrada

- Cisco

```
ipv6 prefix-list BGPIn-IPv6-AS64513 description Prefijos Preferidos del
AS64513
ipv6 prefix-list BGPIn-IPv6-AS64513 seq 10 permit 2001:DB8:300:8000::/49
```

- Juniper

```
policy-statement BGPIn-IPv6-AS64513 {
  term term-1 {
    from {
      route-filter 2001:db8:300::/49 exact;
    }
    then accept;
  }
  term implicit-deny {
    then reject;
  }
}
```

282

Este ejemplo muestra el establecimiento de una política de entrada a través de un *prefix-list* en un router Cisco y un *route-filter* en un router Juniper.

En el ejemplo del router Cisco, el *prefix-list* identifica el segundo /49 del AS 64513 que será recibido.

```
ipv6 prefix-list BGPIn-IPv6-AS64513 seq 10 permit
2001:DB8:300:8000::/49
```

En el ejemplo del router Juniper, el *route-filter* identifica el primer /49 del AS 64513 que será recibido.

```
route-filter 2001:db8:300::/49 exact;
```

Un ejemplo de implementación de estas políticas en una configuración IPv4 sería:

- Cisco: `ip prefix-list BGPIn-IPv6-AS64513 seq 10 permit 10.3.128.0/17`
- Juniper: `route-filter 10.3.0.0/17 exact;`

Establecimiento de filtros

- Política de salida

- Cisco

```
ipv6 prefix-list BGPout-IPv6-AS64512 description Prefijos para AS64512
ipv6 prefix-list BGPout-IPv6-AS64512 seq 10 permit 2001:DB8:21::/48
ipv6 prefix-list BGPout-IPv6-AS64512 seq 20 permit 2001:DB8:21:8000::/49
```

- Juniper

```
policy-statement BGPout-IPv6-AS64511 {
  term term-1 {
    from {
      route-filter 2001:db8:21::/48 exact;
      route-filter 2001:db8:21::/49 exact;
    }
    then accept;
  }
  term implicit-deny {
    then reject;
  }
}
```

283

cgi.br

Este ejemplo muestra el establecimiento de una política de salida a través de un *prefix-list* en un router Cisco y un *route-filter* en un router Juniper.

En el ejemplo del router Cisco, los *prefix-list* indican los prefijos que serán anunciados para el AS64512. En este caso se enviará el prefijo /48 IPv6 y el segundo /49.

En el ejemplo del router Juniper, los *route-filter* indican que para el AS 64511 serán anunciados el prefijo /48 IPv6 y el primer /49.

Un ejemplo de esta política de salida en una configuración IPv4 podría ser el siguiente:

- Cisco

```
ip prefix-list BGPout-IPv4-AS64512 seq 10 permit 172.21.0.0/20
ip prefix-list BGPout-IPv4-AS64512 seq 20 permit 172.21.8.0/21
```

- Juniper

```
route-filter 172.21.0.0/20 exact;
route-filter 172.21.0.0/21 exact;
```

Establecimiento de filtros

- Filtros de protección

- Cisco

```
ipv6 prefix-list IPv6-IPv6-AS64501-all description Todos bloques IPv6
ipv6 prefix-list IPv6-IPv6-AS64501-all seq 10 permit 2001:DB8:21::/48 le 128
```

- Juniper

```
policy-statement IPv6-IPv6-AS64501-all {
  term term-1 {
    from {
      route-filter 2001:db8:21::/48 orlonger;
    }
    then accept;
  }
  term implicit-deny {
    then reject;
  }
}
```

284

Otra política de protección importante es la que impide que los AS reciban anuncios de sus propios prefijos.

Para IPv6 podríamos tener algo similar a lo siguiente:

- Cisco/Quagga

```
ipv6 prefix-list IPv6-IPv6-AS64501-all seq 10 permit 2001:DB8:21::/48 le 128
```

- Juniper

```
route-filter 2001:db8:21::/48 orlonger;
```

Para IPv4 podríamos tener algo similar a lo siguiente:

- Cisco / Quagga

```
ip prefix-list IPv4-AS64501-all seq 10 permit 172.21.0.0/20 le 32
```

- Juniper

```
route-filter 172.21.0.0/20 orlonger;
```

Las reglas que acabamos de ejemplificar indican todos los prefijos posibles dentro de un bloque /48 IPv6 o /20 IPv4. Serán utilizadas en la política de entrada para decir: “no acepto ningún prefijo que sea mío”. Esto ayuda a evitar problemas como el secuestro de prefijos.

Por defecto, el router rechaza todos los prefijos que tengan su ASN para evitar *looping*. Sin embargo, nada impide que otro AS en Internet, intencional o inadvertidamente, genere anuncios del otro prefijo, incluso uno más específico. Si no hay protección el router aceptará y encaminará todo el tráfico interno hacia fuera.

Un ejemplo de este tipo de ocurrencia fue el secuestro del prefijo de YouTube. Por decisión del gobierno de Paquistán, el tráfico de YouTube se debía bloquear para evitar el acceso al trailer de una película anti-islámica. Para cumplir esta orden, el operador Pakistan Telecom generó el anuncio de un prefijo más específico del utilizado por YouTube, con la intención de direccionar todos los accesos al sitio a una página que decía “YouTube was blocked”.

Pero el operador anunció esta nueva ruta a su *upstream provider* (primer error) quien, además de no verificar la nueva ruta (segundo error), la propagó a toda Internet (tercer error). De este modo todo el tráfico de YouTube comenzó a ser direccionado a Paquistán y descartado.

Este fue el caso más famoso, pero los secuestros de bloques de direcciones IP ocurren diariamente, ya sea intencionales o no. Esto se debe a que toda la estructura de Internet y el funcionamiento de BGP fueron definidos en base a una relación de confianza. Esta “inocencia” aun está presente y es lo que mantiene toda la estructura actual.

Hay discusiones sobre los modos de verificar si el AS que está anunciando un determinado prefijo tiene autoridad para hacerlo, similar a lo que hace DNSSec.

Más información:

- <http://www.ietf.org/dyn/wg/charter/idr-charter.html>
- <http://www.ietf.org/dyn/wg/charter/sidr-charter.html>
- <http://www.youtube.com/watch?v=IzLPKuAOe50>
- <http://www.wired.com/threatlevel/2008/02/pakistans-accid/>

Establecimiento de filtros

- Filtros de protección

- Cisco

```
ipv6 prefix-list IPv6-block-deny description Prefixos Gerais Bloqueados
ipv6 prefix-list IPv6-block-deny seq 10 permit 0::/0
ipv6 prefix-list IPv6-block-deny seq 20 permit 0000::/8 le 128
ipv6 prefix-list IPv6-block-deny seq 30 permit 3ffe::/16 le 128
ipv6 prefix-list IPv6-block-deny seq 40 permit 2001:db8::/32 le 128
ipv6 prefix-list IPv6-block-deny seq 50 permit 2001::/33 le 128
ipv6 prefix-list IPv6-block-deny seq 60 permit 2002::/17 le 128
ipv6 prefix-list IPv6-block-deny seq 70 permit fe00::/9 le 128
ipv6 prefix-list IPv6-block-deny seq 80 permit ff00::/8 le 128
```

286

También se pueden agregar *prefix-list* de protección. Observe el siguiente ejemplo en un router Cisco:

```
ipv6 prefix-list IPv6-block-deny seq 10 permit 0::/0
ipv6 prefix-list IPv6-block-deny seq 20 permit 0000::/8 le 128
ipv6 prefix-list IPv6-block-deny seq 30 permit 3ffe::/16 le 128
ipv6 prefix-list IPv6-block-deny seq 40 permit 2001:db8::/32 le 128
ipv6 prefix-list IPv6-block-deny seq 50 permit 2001::/33 le 128
ipv6 prefix-list IPv6-block-deny seq 60 permit 2002::/17 le 128
ipv6 prefix-list IPv6-block-deny seq 70 permit fe00::/9 le 128
ipv6 prefix-list IPv6-block-deny seq 80 permit ff00::/8 le 128
```

Estos *prefix-list* verifican, respectivamente:

- La ruta *por defecto*;
- El primer prefijo /8;
- Direcciones de la red de pruebas 6bone;
- Direcciones para documentación;
- Direcciones de los túneles Teredo;
- Direcciones de los túneles 6to4;
- Direcciones *link-local* (RFC 5735);
- Direcciones Multicast.

Un ejemplo de aplicación de un *prefix-list* de protección para IPv4 podría ser el siguiente:

```
ip prefix-list IPv4-block-deny seq 10 permit 0.0.0.0/0
ip prefix-list IPv4-block-deny seq 20 permit 0.0.0.0/8
ip prefix-list IPv4-block-deny seq 30 permit 127.0.0.0/8
ip prefix-list IPv4-block-deny seq 40 permit 169.254.0.0/16
ip prefix-list IPv4-block-deny seq 50 permit 192.0.2.0/24
ip prefix-list IPv4-block-deny seq 60 permit 10.0.0.0/8
ip prefix-list IPv4-block-deny seq 60 permit 172.16.0.0/12
ip prefix-list IPv4-block-deny seq 80 permit 192.168.0.0/16
```

Estos *prefix-list* verifican, respectivamente:

- La ruta *por defecto*;
- El primer prefijo /8;
- La dirección de *loopback*;
- Direcciones *link-local* (RFC 5735);
- Direcciones de la TEST-NET-1 (RFC 5737);
- Direcciones privadas (RFC 1918).

Establecimiento de filtros

- Filtros de protección

- Juniper

```

policy-statement IPv6-block-deny {
  term term-1 {
    from {
      route-filter 0::/0 exact;
      route-filter 0000::/8 orlonger;
      route-filter 3ffe::/16 orlonger;
      route-filter 2001:db8::/32 orlonger;
      route-filter 2001::/32 longer;
      route-filter 2002::/16 longer;
      route-filter fe00::/9 orlonger;
      route-filter ff00::/8 orlonger;
    }
    then accept;
  }
  term implicit-deny {
    then reject;
  }
}

```

288

También se pueden agregar *route-filter* de protección. Observe el siguiente ejemplo en un router Juniper:

```

route-filter 0::/0 exact;
route-filter 0000::/8 orlonger;
route-filter 3ffe::/16 orlonger;
route-filter 2001:db8::/32 orlonger;
route-filter 2001::/32 longer;
route-filter 2002::/16 longer;
route-filter fe00::/9 orlonger;
route-filter ff00::/8 orlonger;

```

Estos *route-filters* realizan las mismas verificaciones presentadas en el ejemplo anterior para routers Cisco.

Un ejemplo de aplicación de un *prefix-list* de protección para IPv4 podría ser el siguiente:

```

route-filter 0.0.0.0/0 exact;
route-filter 0.0.0.0/8 exact;
route-filter 127.0.0.0/8 exact;
route-filter 169.254.0.0/16 exact;
route-filter 192.0.2.0/24 exact;
route-filter 10.0.0.0/8 exact;
route-filter 172.16.0.0/12 exact;
route-filter 192.168.0.0/16 exact;

```

Más información:

- <http://www.space.net/~gert/RIPE/ipv6-filters.html>

Estabelecendo Filtros

- Filtros de permissão

- Cisco

```
ipv6 prefix-list IPv6-block-permit description Prefixos Gerais Permitidos
ipv6 prefix-list IPv6-block-permit seq 10 permit 2000::/3 le 48
```

- Juniper

```
policy-statement IPv6-block-permit {
  term term-1 {
    from {
      route-filter 2000::/3 prefix-length-range /3-/48
    }
    then accept;
  }
  term implicit-deny {
    then reject;
  }
}
```

289

Como se ha visto en el módulo de Seguridad IPv6, el modo de filtrar las direcciones bogons en IPv6 es diferente de la forma de hacer en IPv4. En el IPv6 es más fácil de liberar el rango de direcciones asignadas y bloquear el resto.

Los ejemplos de `prefix-list` y `policy-statement` que se ve abajo muestran una manera flexible de liberar las direcciones IPv6 disponibles para asignación. Una forma más restringida de hacer este mismo tipo de filtro es permitir uno a uno los rangos de direcciones ya asignados a los RIRs. Estos rangos se pueden encontrar en:

- <http://www.iana.org/assignments/ipv6-unicast-address-assignments/ipv6-unicast-address-assignments.xhtml>

Observe un ejemplo de un router Cisco:

```
ipv6 prefix-list IPv6-block-permit description Prefixos Gerais Permitidos
ipv6 prefix-list IPv6-block-permit seq 10 permit 2000::/3 le 48
```

Observe un ejemplo de un router Juniper:

```
policy-statement IPv6-block-permit {
  term term-1 {
    from {
      route-filter 2000::/3 prefix-length-range /3-/48
    }
    then accept;
  }
  term implicit-deny {
    then reject;
  }
}
```

Aplicación de filtros

- Política de entrada

- Cisco

```
route-map BGPIn-IPv6-AS64512 deny 10
match ipv6 address prefix-list IPv6-AS64501-all
!
route-map BGPIn-IPv6-AS64512 deny 20
match ipv6 address prefix-list IPv6-block-deny
!
route-map BGPIn-IPv6-AS64512 permit 30
match ipv6 address prefix-list BGPIn-IPv6-AS64513
match as-path 300
set local-preference 150
!
route-map BGPIn-IPv6-AS64512 permit 40
match ipv6 address prefix-list IPv6-block-permit
```

290

Después de establecer las condiciones de los filtros usando *prefix-list* en los routers Cisco, éstas se deben aplicar a través de *route-maps*.

Por ejemplo:

```
route-map BGPIn-IPv6-AS64512 deny 10
match ip address prefix-list IPv6-AS64501-all
!
route-map BGPIn-IPv6-AS64512 deny 20
match ip address prefix-list IPv6-block-deny
!
route-map BGPIn-IPv6-AS64512 permit 30
match ip address prefix-list BGPIn-IPv6-AS64513
match as-path 300
set local-preference 150
!
route-map BGPIn-IPv6-AS64512 permit 40
match as-path 32
```

Este es el *route-map* de entrada. Tiene un nombre que identifica su función, ya que es posible tener varios *route-maps*. En este ejemplo, indica cómo tratar lo que se reciba del AS 64512. Hay 4 reglas: 10, 20, 30 y 40.

Los *route-maps* trabajan con tests lógicos tipo “y” y “o”; cada prefijo pasa por el *route-map* y si la comparación entre la regla establecida y el prefijo coincide el prefijo es procesado y la comparación finaliza. Si la comparación no coincide se analiza la regla siguiente. Es decir, cada regla es un test de tipo “o”.

La tercera regla tiene dos tests de comparación. Dos o más tests en la misma regla equivalen a una condición tipo “y”. De estem modo, la tercera regla dice que la primera y segunda línea deben coincidir.

El modo de funcionamiento de los *route-map* no depende del hecho de que la configuración sea de una sesión IPv4 o IPv6.

La primera regla descarta todos los prefijos que coinciden con el que se estableció en el `prefix-list IPv6-AS64501-all`, que representa a todos los prefijos del propio AS. Es la regla que protege contra el secuestro de bloques.

La segunda regla descarta los bloques de uso privado especificados en `prefix-list IPv6-block-deny`.

La tercera regla es donde se modificará la *LOCAL_PREFERENCE*. Verifica de cuál AS viene el prefijo (`as-path 300`) y si es el prefijo esperado (`prefix-list BGPin-IPv6-AS64513`). Si coincide con las dos condiciones, el valor de *LOCAL_PREFERENCE* se aumenta a 150. El valor por defecto de *LOCAL_PREFERENCE* es 100 y cuanto mayor sea su valor mayor será la preferencia.

La última regla especificada en el `prefix-list IPv6-block-permit`, permite recibir anuncios de prefijos en el rango reservado por IANA para asignación, el `2000::/3`. Ella permite anuncios de prefijos hasta /48, el tamaño normalmente acepto por las operadoras.

Si el anuncio recibido no es validado por los *route-maps*, él será descartado, ya que los routers Cisco tienen un "*deny*" como última regla implícita.

Aplicación de filtros

- Política de entrada

- Juniper

```
policy-statement nh-BGPIn-IPv6-AS64511 {
  term term-1 {
    from policy IPv6-AS64501-all;
    then reject;
  }
  term term-2 {
    from policy IPv6-block-deny;
    then reject;
  }
  term term-3 {
    from {
      as-path AS64513;
      policy BGPIn-IPv6-AS64513;
    }
    then {
      local-preference 150;
    }
  }
  term accept {
    then accept;
  }
}
```

292

La política implementada en el router Juniper es similar a la del router Cisco presentada anteriormente. La principal diferencia es que las políticas se aplican a los anuncios recibidos del AS 64511.

Por ejemplo:

```
policy-statement nh-BGPIn-IPv6-AS64511 {
  term term-1 {
    from policy IPv6-AS64501-all;
    then reject;
  }
  term term-2 {
    from policy IPv6-block-deny;
    then reject;
  }
  term term-3 {
    from {
      as-path AS64513-IPv4;
      policy BGPIn-IPv6-AS64513;
    }
    then {
      local-preference 150;
    }
  }
  term accept {
    then accept;
  }
}
```

Aplicación de filtros

- Política de salida

- Cisco

```
route-map BGPout-IPv6-AS64512 permit 10
match ipv6 address prefix-list BGPout-IPv6-AS64512
```

- Juniper

```
policy-statement nh-BGPout-IPv6-AS64511 {
  term term-1 {
    from policy BGPout-IPv6-AS64511;
    then accept;
  }
  term implicit-deny {
    then reject;
  }
}
```

293

De la misma manera que se aplicaron las políticas de entrada, podemos aplicar las políticas de salida a través de un route-map y de un policy-statement en el router Juniper. Las prefixes-list y los route-filters aplicados en las políticas de salida se pueden analizar en el slide 287.

Verificación de Vecinos BGP

- Mostrar todos los vecinos BGP IPv4:
 - `show ip bgp summary` (Cisco / Quagga)
 - `show bgp summary` (Juniper)
- Mostrar todos los vecinos BGP de ambas familias:
 - `show bgp ipv4 unicast summary` (Cisco / Quagga)
 - `show bgp ipv6 unicast summary` (Cisco / Quagga)
 - `show bgp all summary` (Cisco / Quagga)

294

Los siguientes comandos listan una serie de informaciones de estado de la tabla BGP:

Mostrar todos los vecinos BGP IPv4:

```
show ip bgp summary (Cisco / Quagga)
show bgp summary (Juniper)
```

Mostrar todos los vecinos BGP de ambas familias:

```
show bgp ipv4 unicast summary (Cisco / Quagga)
show bgp ipv6 unicast summary (Cisco / Quagga)
show bgp all summary (Cisco / Quagga)
```

Estos datos permiten detectar una serie de problemas de la sesión BGP.

Verificación de Vecinos BGP

• Cisco

```
router-R13#show bgp ipv6 unicast summary
BGP router identifier 172.21.15.253, local AS number 64501
BGP table version is 45, main routing table version 45
28 network entries using 4368 bytes of memory
54 path entries using 4104 bytes of memory
45/17 BGP path/bestpath attribute entries using 7560 bytes of memory
34 BGP AS-PATH entries using 848 bytes of memory
0 BGP route-map cache entries using 0 bytes of memory
0 BGP filter-list cache entries using 0 bytes of memory
Bitfield cache entries: current 2 (at peak 3) using 64 bytes of memory
BGP using 16944 total bytes of memory
26 received paths for inbound soft reconfiguration
BGP activity 49/1 prefixes, 96/21 paths, scan interval 60 secs
```

Neighbor	V	AS	MsgRcvd	MsgSent	TblVer	InQ	OutQ	Up/Down	State/PfxRcd
2001:DB8:21::254	4	64501	1867	1856	45	0	0	1w0d	Active
2001:DB8:21::255	4	64501	4136	3642	45	0	0	1d07h	26
2001:DB8:20::255	4	64512	1896	1876	45	0	0	1d07h	0

295

cgi.br

En este ejemplo podemos observar los siguientes datos sobre los vecinos de las sesiones BGP IPv6 en un router Cisco:

- Neighbor – IP del vecino en que se estableció la sesión BGP;
- V – versión de BGP;
- AS – ASN del vecino;
- MsgRcvd – cantidad de mensajes recibidos del vecino;
- MsgSent – cantidad de mensajes enviados al vecino;
 - estos dos últimos campos generalmente no son verificados, sin embargo son importantes. Muchos cambios en los valores de estos campos pueden indicar un problema. Recibir muchos mensajes, si la tabla se actualiza con frecuencia, puede indicar una fluctuación grande con su vecino.
- TblVer – versión de la tabla;
- InQ – cola de entrada de paquetes;
- OutQ – cola de salida de paquetes;
- Up/Down - momento del último cambio de estado;
- State/PfxRcd – indica el estado actual o el número de prefijos aprendidos. Recuerde que a pesar de la existencia de estados tales como Active y Established que aparentemente indican que la sesión está OK, éstos solo representan estados intermedios de la conexión. La sesión solo estará plenamente establecida cuando haya una indicación de cuántos prefijos fueron aprendidos.

Observe la salida del mismo comando, pero ahora para visualizar las informaciones sobre los vecinos BGP IPv4 en un router Cisco:

```
router-R13#show bgp ipv4 unicast summary
BGP router identifier 172.21.15.253, local AS number 64501
BGP table version is 80, main routing table version 80
31 network entries using 4092 bytes of memory
52 path entries using 2704 bytes of memory
33/22 BGP path/bestpath attribute entries using 5544 bytes of memory
28 BGP AS-PATH entries using 672 bytes of memory
0 BGP route-map cache entries using 0 bytes of memory
0 BGP filter-list cache entries using 0 bytes of memory
Bitfield cache entries: current 3 (at peak 4) using 96 bytes of memory
BGP using 13108 total bytes of memory
21 received paths for inbound soft reconfiguration
BGP activity 99/40 prefixes, 185/84 paths, scan interval 60 secs
```

Neighbor	V	AS	MsgRcvd	MsgSent	TblVer	InQ	OutQ	Up/Down	State/PfxRcd
10.2.255.255	4	64512	10578	10474	80	0	0	1w0d	Active
172.21.15.254	4	64501	10544	10490	80	0	0	1w0d	0
172.21.15.255	4	64501	10572	10490	80	0	0	1w0d	21

Verificación de Vecinos BGP

• Juniper

```

juniper@R11> show bgp summary
Groups: 4 Peers: 5 Down peers: 1
Table      Tot Paths  Act Paths Suppressed  History Damp State  Pending
inet.0     19         17         0           0         0     0       0
inet6.0    56         27         0           0         0     0       0
Peer      AS      InPkt  OutPkt  OutQ  Flaps  Last Up/Dwn  State|#Active/Received/Accepted/Damped
10.1.8.1   64511   3785   4127    0     0    1d 7:26:53  17/17/17/0
172.28.15.252 64508  3776   4135    0     0    1d 7:26:38  0/2/2/0
172.28.15.254 64508  3775   4136    0     0    1d 7:26:46  Connect
2001:db8:28::252 64508  3794   4147    0     0    1d 7:26:40  Establ inet6.0: 0/29/29/0
2001:db8:28::254 64508  3775   4149    0     0    1d 7:26:46  Establ inet6.0: 0/0/0/0
2001:db8:10::1 64511  3810   4128    0     0    1d 7:26:57  Establ inet6.0: 27/27/27/0

```

297

cgi.br

En este ejemplo podemos observar los siguientes datos sobre los vecinos de las sesiones BGP en un router Juniper:

- Groups - número de grupos BGP;
- Peers – número de vecinos BGP;
- Down peers – número de vecinos BGP desconectados;
- Table – nombre de la tabla de enrutamiento;
- Tot Paths – número total de caminos;
- Act Paths – número de rutas activas;
- Suppressed - número de rutas actualmente inactivas. Estas rutas no aparecen en la tabla de enrutamiento y no son exportadas por los protocolos de enrutamiento;
- History - número de rutas retiradas almacenadas localmente para mantener el control histórico de la inestabilidad;
- Damp State – número de rutas con un valor de mérito mayor que cero, pero que continúan activas porque el valor no llegó al límite en que ocurre la retirada;
- Pending – rutas que están siendo procesadas por la política de importación del BGP;
- Peer – dirección de cada vecino BGP;
- AS – ASN del vecino;
- InPkt – número de paquetes recibidos del vecino;
- OutPkt – número de paquetes enviados al vecino;
- OutQ – cola de salida de paquetes;
- Flaps – número de veces que la sesión BGP fue interrumpida y se restableció;
- Last Up/Down – última vez que ocurrió un cambio de estado;
- State|#Active/Received/Accepted/Damped – indica el estado actual o el número de prefijos aprendidos. Si la sesión no ha sido establecida, este campo muestra el estado actual de la sesión: Active, Connect, o Idle. Si la sesión ha sido establecida, el campo indica el número de rutas activas, recibidas, aceptadas o inestables.

Observe que el router Juniper muestra a la salida del propio comando las informaciones sobre las sesiones IPv4 e IPv6.

Looking Glass

- Es importante verificar a través de *Looking Glasses* remotos cómo los operadores y toda la Internet reciben los anuncios del AS.
- Cisco

```
bgpd-R01> show bgp regexp _64501$
BGP table version is 0, local router ID is 10.3.255.255
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale, R Removed
Origin codes: i - IGP, e - EGP, ? - incomplete

   Network        Next Hop        Metric LocPrf Weight Path
*  2001:db8:21::/48 2001:db8:300:11::2          0 64511 64501 i
*>                2001:db8:300:12::2          0 64512 64501 i
*  2001:db8:21::/49 2001:db8:300:11::2          0 64511 64512 64501 i
*>                2001:db8:300:12::2          0 64512 64501 i
*  2001:db8:21:8000::/49
                    2001:db8:300:12::2          0 64512 64511 64501 i
*>                2001:db8:300:11::2          0 64511 64501 i
Total number of prefixes 3
```

298

Es importante verificar a través de *Looking Glasses* remotos cómo los operadores y toda la Internet reciben los anuncios del AS. De esta manera también es posible verificar si las políticas de enrutamiento han sido bien aplicadas.

Con el *Looking Glass* es posible consultar cómo los prefijos de un AS, tanto IPv4 como IPv6, están siendo propagados por Internet, es decir, cómo los AS pueden comunicarse con su red.

Por ejemplo:

`show bgp regexp _64501$` (Cisco / Quagga)

En este ejemplo podemos observar cómo en un router Cisco fue aprendido cada prefijo anunciado por el AS 64501. En esta expresión regular el caracter “\$” significa que es el origen (inicio de la línea), y el caracter “_” significa espacio, es decir, el comando se aplica para todo prefijo que tenga al AS 64501 como origen en el AS-PATH.

En respuesta a la consulta podemos ver el balanceo de tráfico (un bloque /48 se dividió en dos prefijos /49 permitiendo que la mitad del tráfico salga por un *enlace* y la otra mitad por otro) y también la redundancia de rutas (además de los prefijos /49, el prefijo /48 también está siendo anunciado por los dos *enlaces*.)

Looking Glass

- Juniper

```
juniper@R11> show route table inet6.0 aspath-regex .64513$

inet6.0: 59 destinations, 84 routes (59 active, 0 holddown, 0 hidden)
+ = Active Route, - = Last Active, * = Both

2001:db8:300::/48 *[BGP/170] 01:44:51, localpref 100
                   AS path: 64511 64513 I
                   > to 2001:db8:100:1::1 via ge-0/0/0.2105
                   [BGP/170] 01:44:13, MED 0, localpref 100, from 2001:db8:21:ffff::252
                   AS path: 64512 64513 I
                   > to fe80::224:97ff:fec1:c8bd via ge-0/0/0.2101
2001:db8:300::/49 *[BGP/170] 01:44:13, MED 0, localpref 150, from 2001:db8:21:ffff::252
                   AS path: 64512 64513 I
                   > to fe80::224:97ff:fec1:c8bd via ge-0/0/0.2101
                   [BGP/170] 01:44:51, localpref 100
                   AS path: 64511 64513 I
                   > to 2001:db8:100:1::1 via ge-0/0/0.2105
2001:db8:300:8000::/49
                   *[BGP/170] 01:44:51, localpref 150
                   AS path: 64511 64513 I
                   > to 2001:db8:100:1::1 via ge-0/0/0.2105
```

299

En este ejemplo podemos observar, en un router Juniper, cómo los prefijos anunciados por el AS64513 fueron aprendidos por el AS64501. La expresión regular utilizada es similar a la descrita en el ejemplo anterior utilizando routers Cisco.

Igual que en el ejemplo anterior, en respuesta a la consulta podemos observar el balanceo del tráfico y la redundancia de rutas, además de datos tales como: mejor ruta, indicada por el asterisco; camino hasta el destino (*AS path*); próximo salto; e interfaz de salida del paquete.

IPv6.br

La nueva generación del
Protocolo de Internet



Planificación

Módulo 10

301

cgi.br

Ahora que ya hemos aprendido un poco sobre el funcionamiento de los nuevos recursos del protocolo IPv6, en este módulo discutiremos algunos conceptos importantes que deben ser analizados a la hora de implementar IPv6 en las grandes redes. Debemos comprender todos los aspectos que implica la implementación de IPv6 principalmente en las redes empresariales, como el análisis de costos, el recambio de equipos y la capacitación.

Planificación

- La decisión sobre la adopción del protocolo IPv6 genera muchas preguntas.
 - ¿IPv6 es realmente necesario?
 - ¿Cuándo será necesario tener IPv6?
 - ¿Hay alternativas viables al uso de IPv6?
 - ¿La transición se debe realizar de una vez o gradualmente?
 - ¿Cómo hacer para que las aplicaciones y servicios sean compatibles con el nuevo protocolo?
 - ¿Cómo aprovechar las nuevas funcionalidades de IPv6?
 - Además de la seguridad, ¿qué aspectos deben ser evaluados?
 - ¿Cómo planificar para esta transición?
 - ¿Cuál es el costo de implementación?

302

egi.br

Conceptos importantes:

- El cambio de protocolo un cambio de naturaleza estructural;
- Este cambio no se producirá simplemente porque el protocolo IPv6 presenta mejoras respecto de su predecesor;
- La implementación de IPv6 es necesaria e inevitable;
- El agotamiento de las direcciones IPv4 no acabará con la Internet ni hará que deje de funcionar;
- Pero habrá una disminución de la tasa de crecimiento de la red y dificultadas para el desarrollo de nuevas aplicaciones;
- Todos estos problemas se pueden evitar adoptando el protocolo IPv6 antes del agotamiento del espacio IPv4;
- La implementación de IPv6 no será algo rápido;
- No hay una fecha definitiva para el cambio de protocolo;
- La migración de IPv4 a IPv6 se producirá de forma gradual, mientras IPv4 continúa funcionando;
- Es importante que desde ya las redes estén preparadas para el nuevo protocolo. Cuanto más pronto se comprenda el tema y se planifique la implementación, menores serán los gastos del proceso.

Primer paso: Capacitación

- Es importante que tanto los técnicos como los administradores de redes busquen adquirir conocimiento sobre esta nueva tecnología;
 - Cursos;
 - Libros;
 - Sitios web;
 - Documentos técnicos;
 - Foros;
 - Eventos.
- Este primer paso será esencial para la elaboración de las siguientes fases.

303

cgi.br

RIPE NCC ofrece algunos videos sobre “casos IPv6” que pueden servir como una muy interesante fuente de conocimiento. Estos videos están disponibles en el canal ripencc de YouTube: www.youtube.com/ripencc, y algunos están subtítulados:

- Randy Bush (Internet Initiative Japan Inc.) - <http://www.youtube.com/watch?v=Qh3i6lDqWBM>
- Lorenzo Colitti (Google) - <http://www.youtube.com/watch?v=vFwStbTpr6E>
- Marco Hogewoning (Dutch ISP XS4ALL) - <http://www.youtube.com/watch?v=f3WcWBIQ11A>
- Andy Davidson (NetSumo ISP Consultancy) - <http://www.youtube.com/watch?v=QCcigLJJbvU>
- David Freedman (Claranet) - <http://www.youtube.com/watch?v=HQtbz1ahRxE>

El impacto de IPv6

- De qué manera IPv6 puede afectar los negocios:
 - Nuevas aplicaciones;
 - Nuevas oportunidades;
 - Nuevos servicios;
- Cómo obtener conexión;
- Qué tipo de conexión ofrecer a los clientes;
 - IPv6 nativo;
 - Túneles.
- Qué servicios internos se migrarán inicialmente;
- Entender estos aspectos es fundamental para optimizar el retorno sobre las inversiones.

304

El impacto de IPv6

- Minimizar los costos de implementación.
 - Costo relativo en un ISP*:
 - Con equipos (15%)
 - Routers - Intermedio;
 - Firewalls - Intermedio.
 - Con software (15%)
 - Software de administración y monitoreo de redes - Alto;
 - SOs - Intermedio.
 - Mano de obra (70%)
 - Investigación y desarrollo - Bajo;
 - Capacitación - Alto;
 - Implementación - Alto;
 - Mantenimiento – Intermedio/Alto;
 - Problemas de interoperabilidad - Intermedio/Alto.

*Fuente: U.S. DEPARTMENT OF COMMERCE

305

Ejercicio Ingenieros x Gerentes

Gerentes/Ejecutivos/Directores

X

Técnicos/Ingenieros

- ¿Usted puede convencer a los ejecutivos de su empresa?
- Vamos a trabajar en grupos...
 - 10 min – preparación
 - 10 min – presentación y debate

306

egi.br

Ejercicio

Ingenieros x Gerentes

- Gerentes:
 - No quieren invertir en IPv6 en este momento
- Técnicos:
 - Quieren convencer a los gerentes que es necesario actuar ahora
- Aspectos a considerar:
 - Capacidad del hardware
 - Prioridades del negocio
 - Conocimiento existente / Capacitación
 - Clientes
 - Legislación
 - Costo
 - *Timing*
 - Intercambio de tráfico

307

Este ejercicio se utilizó en la capacitación sobre IPv6 que RIPE ofreció a los proveedores europeos. Las siguientes son algunas de las respuestas que se presentaron:

Técnicos:

- ¡Un cliente lo pidió!
- Dispositivos móviles
- Nuestros equipos ya lo soportan
- Si lo hacemos ahora podemos invertir gradualmente
- Solicitar la asignación en este momento es una inversión, mucho espacio, ningún costo
- Hay diferentes oportunidades de capacitación
- V6 es la única opción para el crecimiento
- V6 es fácil, no requiere demasiados recursos
- Compare el precio de adoptar IPv6 ahora con el costo de hacerlo más adelante

Gerentes:

- Va a costar mucho, vamos a necesitar 20% más de presupuesto
- No tenemos clientes para esto
- Antes debemos auditar los equipos
- Necesitamos un plan, con diferentes fases
- Los equipos pueden tener v6 pero no hay paridad en las funcionalidades
- Tenemos equipos y software obsoletos
- Las empresas que están preocupadas por la seguridad no quieren IPv6 ahora porque el hardware y el software aun no están maduros
- ¿Por qué debemos cambiar?
- ¿Cómo me puede garantizar que v4 realmente se va a agotar?

Escenario: No hacer nada

- Ningún problema en los próximos años
- Con el paso del tiempo algunas personas no podrán utilizar sus servicios
- Ningún costo extra
 - ¡Hasta que nos demos de cabeza contra el muro!
- Altos costos para una implementación rápida
- Tiempos de planificación más cortos implican más errores...

Escenario: Hacer todo ahora

- Puede que sea necesario cambiar el hardware
- Gran inversión de tiempo y otros recursos
- Sin retorno inmediato
- Altos costos para una implementación rápida
- Una planificación rápida significa mayor probabilidad de error...

Escenario: Empzar ahora, hacerlo en etapas

- Definir metas y plazos a cumplir.
- Identificar qué áreas y servicios serán afectados.
- Procedimientos de compra
 - Paridad de funcionalidades
- Estudie su hardware y software
 - Aplicaciones o sistemas que no se actualizarán;
 - Servicios críticos.
- Realice pruebas
- Un servicio a la vez:
 - Primera fase
 - Core
 - Clientes
- Prepárese para deshabilitar IPv4

Primera fase

- Desenvolver um plano de direcionamento
- Obter as direções
- Anunciá-las
- Web
- DNS autoritativo
- Servidores de correio eletrônico
- etc.

Obtener un prefijo IPv6

- Todos los RIR ya están distribuyendo direcciones IPv6 en sus regiones.
- Complete el formulario en:
 - <http://registro.br/info/pedido-form.txt>
- Enviar por email: numeracao-pedido@registro.br
- Recibirá un recibo o un mensaje de error sobre los datos incluidos en el formulario
- Quien tiene IPv4 seguramente justifica IPv6
- Gratis, por ahora
- 2 semanas entre análisis y aprobación
- Dudas: numeracao@registro.br

Lo que no se debe hacer

- No separar las funcionalidades v6 de v4
- No hacer todo de una vez
- No designar un "gurú de IPv6" para su organización
 - ¿Tiene un especialista en IPv4?
- No considerar a IPv6 como un producto
 - El producto es Internet, o el acceso a Internet o sus contenidos

Consideraciones

- IPv4 ya no es sinónimo de Internet
- Evitar el problema no lo hará desaparecer
- ¿Cuánto está dispuesto a gastar ahora para ahorrar dinero después?
- Solo IPv6 permitirá el crecimiento continuo de la red

¡Empiece ahora!

BIBLIOGRAFIA

- Migrating to IPv6 : A Practical Guide to Implementing IPv6 in Mobile and Fixed Networks .
Autor: Marc Blanchet.
- 6 Net: An IPv6 Deployment Guide.
Autor: Martin Dunmore.
- IPv6 Essentials.
Autor: Silvia Hagem.
- Global IPv6 Strategies.
Autores: Patric Grossetete; Ciprian popovicius; Fred Wetting.
- Planning and Accomplishing the IPv6 Integration: Lessons Learned from a Global Construction and Project-Management Company .
Autor: Cisco Public Information.
- Technical and Economic Assessment of Internet Protocol Version 6.
Autor: U.S. DEPARTMENT OF COMMERCE.
- Introducción a IPv6.
Autor: Roque Gagliano.
- Planificando IPv6.
Autor: Roque Gagliano.
- Deliverable D 6.2.4: Final report on IPv6 management tools, developments and tests.
Autor: 6 Net.
- IPv6 Security: Are You Ready? You Better Be!
Autor: Joe Klein.
- IPv6 and IPv4 Threat Comparison and Best- Practice Evaluation (v1.0)
Autores: Sean Convery; Darrin Miller.
- BGP Routing Table Analysis Reports - <http://bgp.potaroo.net/>
Autores: Tony Bates; Philip Smith; Geoff Huston.
- Measuring IPv6 Deployment
Autores: Geoff Huston; George Michaelson.
- IPv6 at Google.
Autores: Angus Lees; Steinar H. Gunderson.
- Resumo do Barômetro Cisco Banda Larga Brasil 2005-2010
Autores: Mauro Peres; João Paulo Bruder.
- Tracking the IPv6 Migration. Global Insights From the Largest Study to Date on IPv6 Traffic on the Internet.
Autor: Craig Labovitz.
- ICE: Uma solução geral para a travessia de NAT
Autor: José Henrique de Oliveira Varanda

BIBLIOGRAFIA

- TIC Domicílios e Usuários Total Brasil - <http://cetic.br/usuarios/tic/2009-total-brasil/index.htm>
Autor: CETIC.br.
- IPv6.br - <http://ipv6.br>
Autor: CEPTR0.br.
- IPv6 Deployment and Support - <http://www.6deploy.org>
Autor: 6Deploy.
- World Internet Usage Statistics News and World Population Stats -
<http://www.internetworldstats.com/stats.htm>
Autor: Internet World Stats.
- Barômetro – BRASIL - <http://www.cisco.com/web/BR/barometro/barometro.html>
Autor: Cisco Systems.
- The ISC Domain Survey - <https://www.isc.org/solutions/survey>
Autor: Internet Systems Consortium.
- Number Resource Organization – Statistics - <http://www.nro.net/statistics>
Autor: Number Resource Organization (NRO).
- Routing TCP/IP
Autor: Jeff Doyle, Jennifer DeHaven Carroll
- O Protocolo BGP4 - Parte 3 (Final) - <http://www.rnp.br/newsgen/9907/pgbp4p3.html>
Autor: RNP – Rede Nacional de Ensino e Pesquisa