

UNIVERSIDAD NACIONAL AUTÓNOMA DE NICARAGUA

UNAN – León

Facultad de Ciencias y Tecnología

Departamento de Computación



Propuestas de prácticas de laboratorios de Switching, Routing y Servicios de Red con IPv6 para la asignatura “Despliegue de IPv6” correspondiente a la Electiva X de la carrera de Ingeniería en Telemática del Departamento de Computación de la UNAN-León.

Tesis para optar al título de
INGENIERO EN TELEMÁTICA

Presentado por:

Br. Delia María Jaime Toruño.

Br. Hugo Mariano García Machado.

Br. William Francisco Aguilar Zapata.

Tutor:

MSc. Aldo René Martínez Delgadillo.

León, Junio 2015



Le agradezco:

A Dios, por esta siempre a mi lado permitiéndome llegar a este momento tan especial en mi vida, por los triunfos y los momentos difíciles que me han enseñado a valorarlo cada día más.

A mis padres Apolonia Toruño y Leopoldo Jaime, por la confianza que tuvieron en cada momento de mi vida, por su amor y todo el esfuerzo y sacrificio que pasaron para que yo pudiera culminar mis estudios.

A Mercedes Morales y Carlos Toruño, porque fueron como mis segundos padres que me apoyaron incondicionalmente durante el transcurso de la carrera.

A mis hermanos y familia que siempre pusieron un granito de arena y por los momentos que siempre compartimos juntos.

A mis amigos y amigas por compartir conmigo cada momento especial, y por la amistad incondicional que siempre me brindaron.

A mi tutor Aldo Martínez, por el tiempo que nos dedicó, por el apoyo que nos brindó para la realización de la tesis monográfica y por la paciencia que demostró en el transcurso de este trabajo.

Delia María Jaime Toruño.



En primer lugar quiero agradecerle a Dios, nuestro padre, por la vida que me ha dado, por los momentos lindos que he tenido, por ser la luz que guía mis pasos para seguir adelante y lograr las metas que me propuesto.

Le doy gracias a mi madre Juana Mercedes Machado Martínez y a mi padre Hugo Mariano García Camacho, por ser unos padres ejemplares, maravillosos y dedicados. Por estar siempre a mi lado, de brindarme su amor y sobre todo de guiarme en esta vida para ser un hombre de éxito.

Le doy gracias a mi abuelo Hugo Mariano García Morales y a mi abuela Angélica Camacho Espinoza, por estar siempre en mi vida, demostrarme el gran amor que me han dado, de inculcarme valores, que hasta el día de hoy han sido de gran enseñanza. Hoy más que nunca quiero decirle, de lo orgulloso que estoy de que ustedes sean mis abuelos y yo de ser su nieto.

Le doy gracias al amor de mi vida, Leyla Yesenia Baca Suazo, por ser una mujer excepcional, de apoyarme y brindarme su amor en cada momento de mi vida.

Le doy gracias a mi familia, por la confianza y el apoyo que me han brindado en el transcurso de mi vida.

Le doy gracias a mis compañeros de la tesis monográfica, Delia María Jaime Toruño y William Francisco Aguilar Zapata, por haberme tenido la paciencia necesaria y por motivarme a seguir adelante en los momentos difícil.

Le doy gracia al tutor de la tesis Aldo Rene Martínez Delgadillo, por transmitirnos sus conocimientos acerca del tema y del apoyo brindado en cada etapa de la elaboración de la tesis.

Hugo Mariano García Machado.



A Dios.

Por haberme permitido llegar hasta este punto y haberme dado salud para lograr mis objetivos, además de su infinita bondad y amor.

A mi madre María E. Zapata.

Por haberme brindado su apoyo en todo momento, por sus consejos, sus valores, por la motivación constante que me ha permitido ser una persona de bien, pero más que nada, por su amor.

A mis amigos

Que nos apoyamos mutuamente en nuestra formación profesional y que hasta ahora, seguimos siendo amigos: Hugo García y Delia Jaime. Además de otros amigos como son José Espinoza, Ligia Espinoza, María G. Martínez y Meyling G. Vargas, etc.

William Francisco Aguilar Zapata



CONTENIDO

CAPITULO I: ASPECTOS INTRODUCTORIOS	1
1. Introducción	2
2. Antecedentes	3
3. Planteamiento del Problema	4
4. Justificación	5
4.1 Originalidad.....	5
4.2 Alcance.	5
4.3 Producto.	6
4.4 Impacto.	6
5. Objetivos.....	7
5.1 Objetivo general.....	7
5.2 Objetivos específicos.	7
6. Diseño metodológico	8
6.1 Recolección de Información.....	8
6.2 Selección de las herramientas a implementar	8
6.3 Elaboración y desarrollo de los laboratorios	8
CAPITULO II: DESARROLLO TEÓRICO	12
1. IPv4	13
1.1 Definición	13
1.2 Paquete o Datagrama.....	14
1.3 Direccionamiento.	15
1.4 Clases de direccionamiento.....	16
1.5 Ruteo interno de dominio sin clases (CIDR).	18
1.6 Problemas con IPv4.....	21
1.7 Solución a IPv4.....	22
1.8 ¿Por qué cambiar a IPv6?	23
1.9 Ejercicios de IPv4.	24
2. IPv6	26



2.1 Introducción	26
2.2 Características.....	27
2.3 Datagrama IPv6	29
2.4 Modificación de cabecera IPv4 a IPv6.....	30
2.5 Encabezado de Extensión	32
2.6 Tipos de cabeceras de Extension.....	35
2.7 Ejercicios de análisis de encabezado de extensión IPv6.....	42
2.8 Notación IPv6.....	44
2.9 Direcciones IPv6	47
2.10 Unicast (Identificación Individual).	48
2.11 Tipos de direcciones Unicast.....	48
2.12 Anycast (Identificación Selectiva).	51
2.13 Multicast.....	53
2.14 Plan de direccionamiento.....	55
2.16 Subneting en IPv6.....	56
2.17 Ejercicio de IPv6.	59
3. ICMPv6 (Control Internet Message Protocol version 6)	61
3.1 Formato de mensajes ICMPv6	61
3.2 Mensajes de error ICMPv6	62
3.3 Ejemplo Destino Inalcanzable.....	64
3.4 Mensajes de Información.....	64
3.5 Descubrimiento de Vecinos	65
4. Autoconfiguración	69
4.1 Autoconfiguración “stateless” (sin intervención).	69
4.2 Autoconfiguración “stateful” (predeterminada).....	70
4.3 Ejemplo de autoconfiguración.....	71
5. Enrutamiento con IPv6.	73
5.1 Tipos de protocolos de enrutamiento.....	73
5.2 IGP (Protocolos de Gateway interior)	73
5.3 EGP (Protocolo de pasarela externo).	74



5.4 Características de los protocolos de enrutamiento IGP y EGP	74
5.5 Diferencias entre IGP y EGP	75
5.6 Sistemas Autónomos.....	75
5.7 Enrutamiento Estático.....	76
6. DHCPv6.....	80
6.1 Arquitectura Cliente-servidor.	80
6.2 Identificador Único DHCP (DUID).	81
6.3 Mensajes DHCPv6.....	82
6.4 Agente DHCPv6 relay.....	87
6.5 Autenticación de mensajes DHCPv6	89
6.6 DHCPv6 con estado y sin estado.	90
6.7 Ejemplo de configuración de DHCPv6.....	91
7. RIPng (RIP new Generation)	95
7.1 Generalidades de RIPng.....	95
7.2 Formato de mensaje RIPng	96
7.3 Características de RIPng:.....	97
7.4 Configuración RIPng.....	99
8. EIGRPv6.....	103
8.1 Diferencias entre EIGRP IPv6 e IPv4	104
8.2 Ejemplo de configuración de EIGRPv6.....	105
9. OSPFv3	109
9.1 Tipos de paquetes OSPF.....	110
9.2 Diferencias entre OSPFv2 y OSPFv3	111
9.3 Soporte para múltiples instancia OSPFv3	111
9.4 Seguridad	112
9.4 Tipos de LSA para IPv6	113
9.5 Configuración de OSPFv3 en IPv6.	114
10. Integrated IS-ISv6.....	118
10.1 Vecinos y Adyacencias	119
10.2 Procesos de ISIS – Update Process.....	120



10.3 La Capacidad IPv6 TLV (Tipo 236).....	122
10.4 IS-ISv6 Adyacencia.....	125
10.5 IS-IS Single Topology:	125
10.6 IS-IS multi-topology.....	125
10.7 Configuración Integrated IS-IS.....	126
11. BGP4	127
11.1 Configuración Explícita Peers (pares) BGP	128
11.2 Uso de Claves Compartidas en Sesiones BGP	128
11.3 Aprovechamiento de un Túnel IPsec	129
11.4 Configuración de BGP4	130
12. VLANs (Red de área local virtual).....	135
12.1 Tipos de VLANs	135
12.2 Tipos de puertos	136
12.3 Interconexiones de switch con VLANs y puertos trunk	136
12.4 Configuración de VLANs con IPv6.....	137
13. Frame Relay.....	140
13.1 Beneficios de Frame Relay	140
13.2 Funcionamiento de Frame Relay	140
13.3 Circuitos virtuales.....	141
13.4 DLCI (Data Link Connection Identifier)	141
13.5 Mapa Frame Relay	142
13.6 Encapsulación Frame Relay	144
13.7 Configuración de Frame Relay	145
14. IPv6 Access Control Lists	146
14.1 ACLs de entrada y salida.....	146
14.2 Paquetes de inspección en IPv6.....	148
14.3 Limitaciones ACL en IPv6.....	148
14.4 Configuración de ACL en IPv6.....	149
15. Mecanismos de transición.	151
15.1 DSTM.....	151



15.2 Túneles	153
15.3 Traductores	161
16. Movilidad en IPv6	180
16.1 Escenario de la Movilidad en IPv6	181
16.2 Túnel IP Móvil en la movilidad IPv6	182
16.3 Nueva Cabecera de Extensión IPv6	186
16.4 Ejercicio de análisis de movilidad	191
17. Seguridad en IPv6	196
17.1 IPsec	196
17.2 VPN	204
17.3 VPN sobre IPsec	208
18. Servicios en IPv6	210
18.1 DNS	210
18.2 FTP	216
18.3 HTTP en IPv6.	217
18.3 Secure Shell	219
18.4 WiFi (802.11).	222
19. VOIP siguiente generación de voz en IPv6	226
19.1 Migración de voz IP versión 6	226
19.2 Calidad de servicio	228
19.3 Desempeño IPv6 frente a IPv4	228
19.4 Arquitectura de integración.	229
19.5 Medición del retardo	230
CAPITULO III: DESARROLLO PRÁCTICO	232
ORGANIZACIÓN DE LAS PRÁCTICAS.	233
PRÁCTICA 1: DIRECCIONAMIENTO IPv6 CON RUTAS ESTÁTICAS	235
PRÁCTICA 2: AUTOCONFIGURACIÓN Y DHCPv6	240
PRÁCTICA 3: MECANISMOS DE TRANSICIÓN IPv4 E IPv6	245
PRÁCTICA 4: VLANS ESTÁTICAS Y DINÁMICAS	254
PRÁCTICA 5: FRAME RELAY E INTERVLANS	263
PRÁCTICA 6: CONTROL DE LISTAS DE ACCESO.	272



PRÁCTICA 7: PROTOCOLO DE ENRUTAMIENTO INTERNO (RIPng, EIGRPv6, OSPFv3). ...	280
PRÁCTICA 8: PROTOCOLOS DE ENRUTAMIENTO INTERNOS y EXTERNOS.	290
PRÁCTICA 9: REDES VIRTUALES PRIVADAS.....	300
PRÁCTICA 10: VOIP CON IPv6.....	306
PRÁCTICA 11: DNS	323
PRÁCTICA 12: HTTP.....	329
PRÁCTICA 13: INTÉRPRETE DE ÓRDENES SEGURA (SSH).....	335
PRÁCTICA 14: MISCELANEA IPv6.....	342
CAPITULO IV: ASPECTOS FINALES	365
1. CONCLUSIONES.....	366
2. RECOMENDACIONES	367
3. BIBLIOGRAFÍA.....	368



ÍNDICE DE FIGURAS

Figura 1: Etapas del trabajo.....	8
Figura 2: Simuladores usados.....	8
Figura 3: Comunicación con diferentes sistemas de internet.....	13
Figura 4: Datagrama IPv4.....	14
Figura 5: Ejemplo de dirección IPv4.....	15
Figura 6: Formas de direccionar una dirección IPv4.....	16
Figura 7: Distribución de bits de red y host en las diferentes clases de IPv4.....	17
Figura 8: Clase de direccionamiento IPv4.....	17
Figura 9: Características de IPv6.....	27
Figura 10: Forma general de un datagrama de IPv6.....	29
Figura 11: Formato de Encabezado de IPv6.....	29
Figura 12: Campos eliminado en IPv6.....	30
Figura 13: Modificación de campos en IPv6.....	31
Figura 14: Campos mantenido en IPv6.....	31
Figura 15: Modificación entre la cabecera IPv4 e IPv6.....	31
Figura 16: Encabezado de Extensión IPv6.....	32
Figura 17: Tipo de encabezado de Extensión IPv6.....	33
Figura 18: Orden de cabecera de Extensión IPv6.....	33
Figura 19: Formato de cabecera Opción de Salto a Salto y Opción Destino.....	34
Figura 20: Formato Cabecera de Extensión Salto a Salto.....	36
Figura 21: Formato de cabecera de Extensión de Opción Destino.....	37
Figura 22: Formato cabecera de Extensión de Enrutamiento.....	37
Figura 23: Formato cabecera de Extensión de Fragmentación.....	38
Figura 24: Zona divisible vs indivisible de la fragmentación.....	39
Figura 25: Formato cabecera de Extensión de Fragmentación.....	40
Figura 26: Cabecera ESP.....	41
Figura 27: Modo Transporte.....	41
Figura 28: Modo Túnel.....	41
Figura 29: Notación IPv6.....	44
Figura 30: Campos con ceros al inicio y ceros totales.....	44
Figura 31: Campos sucesivos de ceros.....	45
Figura 32: Prefijo IPv6.....	46
Figura 33: Tipo de direcciones IPv6.....	47
Figura 34: Direcciones Unicast Globales Agregables.....	49



Figura 35: Estructura Uso Local.	50
Figura 36: Dirección Link-local	50
Figura 37: Dirección Site-local.....	51
Figura 38: Direcciones IPv6 con direcciones IPv4.....	51
Figura 39: Dirección IPv6 mapeada a IPv4.	51
Figura 40: Anycast con indicador de interfaz igual a cero	52
Figura 41: Ejemplo de direcciones Anycast.....	53
Figura 42: Multicast	53
Figura 43: Direcciones obligatorias en un host IPv6.....	55
Figura 44: Direcciones obligatorias en un router IPv6.....	56
Figura 45: Topología de empresa.....	57
Figura 46: Descomposición de dirección IPv6 de la Empresa.....	57
Figura 47: Separación de los bit de red y bit de host.....	58
Figura 48: Combinación de binario requeridos para subnetear dirección IPv6	57
Figura 49: Esquema final de la empresa.	58
Figura 50: Mensajes de error de ICMPv6.....	63
Figura 51: Ejemplo Destino Inalcanzable	64
Figura 52: Mensajes de Información.	64
Figura 53: Formato de mensajes de información.	65
Figura 54: Mensaje de solicitud del routers.	66
Figura 55: Mensaje de anuncio del routers	66
Figura 56: Mensaje de solicitud vecino.....	67
Figura 57: Mensaje de anuncio de vecino.	68
Figura 58: Autoconfiguración stateless.....	69
Figura 59: Autoconfiguración stateful.	70
Figura 60: Topología de autoconfiguración.	71
Figura 61: Tipo de enrutamiento IPv6	73
Figura 62: Comparación entre IGP y EGP	75
Figura 63: Evolución de los protocolos.....	76
Figura 64: Configuración de ruta estática directamente conectada.....	76
Figura 65: Ruta estática IPv6 totalmente especificada.....	77
Figura 66: Observación de ruta estática IPv6 recursiva	78
Figura 67: Ruta estática IPv6 flotante.....	78
Figura 68: Ruta estática IPv6 por defecto	79
Figura 69: DHCPv6, asignación rápida con dos mensajes.....	84
Figura 70: DHCPv6, asignación con 4 mensajes.	85
Figura 71: DHCPv6, mensaje Renew.....	86



Figura 72: DHCPv6, mensaje Rebind.....	86
Figura 73: Aplicación típica de DHCPv6 Relay Agent.	87
Figura 74: Proceso de funcionamiento de un DHCPv6 relay agent.	88
Figura 75: Configuración automática sin estado (I)	90
Figura 76: Configuración automática sin estado (II).	91
Figura 77: Topología de DHCPv6.....	92
Figura 78: Formato de mensajes RIPng.....	96
Figura 79: Ruteo RTE.....	96
Figura 80: Próximo salto RTE.....	97
Figura 81: Topología de red con RIPng.....	100
Figura 82: Topología de EGRP	105
Figura 83: Topología de OSPv3	115
Figura 84: IPv6 reachability TLV format.	123
Figura 85: IPv6 interface address TLV format.	123
Figura 86: Protocol TLV.....	124
Figura 87: Topología BGP.....	130
Figura 88: Topología de VLANs.	137
Figura 89: Ejemplo de mapeo de circuitos virtuales.	143
Figura 90: Topología de Frame Relay	145
Figura 91: ACL de entrada y salida.	146
Figura 92: Sentencias para consultar mac en IPv6	147
Figura 93: Topología de ACL.....	149
Figura 94: Entorno DSTM.....	152
Figura 95: Entorno 6to4	154
Figura 96: Dirección 6to4.....	154
Figura 97: Conversión de dirección IPv4 a dirección 6to4.....	155
Figura 98: Comunicación 6to4.....	156
Figura 99: Entorno 6over4	157
Figura 100: Diferentes topologías de red con distintos MTU.	159
Figura 101: Entorno Tunnel Broker	161
Figura 102: SIIT para redes pequeñas IPv6.....	162
Figura 103: SIIT para redes Dual Stack.	162
Figura 104: Entorno NAT-PT	168
Figura 105: Componentes de BIS.	172
Figura 106: Representación de NAT-64	175
Figura 107: Función de DNS64	178
Figura 108: Escenario de movilidad en IPv6.	181



Figura 109: Túnel IP móvil en MIPv6.....	182
Figura 110: Optimización del triángulo. O repuestas directas.	183
Figura 111: Formato de mensaje .ND de anuncio de router.....	185
Figura 112: Cabecera de extensión de IPv6.	187
Figura 113: Tráfico de MN a CN.....	189
Figura 114: Cabecera de extensión MIPv6 tráfico de CN a MN.	190
Figura 115: Caso de estudio de MIPv6.	191
Figura 116: Solucion1 de MIPv6.....	193
Figura 117: Solución del inciso 2 de MIPv6.....	194
Figura 118: Solución del inciso3 MIPv6.....	195
Figura 119: IPSec modo transporte.....	196
Figura 120: IPSec Modo Túnel.	197
Figura 121: Implementación AH en modo túnel y modo transporte.....	198
Figura 122: Implementación ESP en modo túnel y modo transporte.	199
Figura 123: Túnel IPv6 en IPv4.	202
Figura 124: Componentes para conexiones VPN a través de Internet IPv4.....	205
Figura 125: Paquetes IPv6 sobre IPv4 que usan una conexión VPN a través de Internet IPv4....	206
Figura 126: Paquetes IPv6 nativo que usan una conexión VPN a través de Internet IPv4	207
Figura 127: Paquetes IPv6 sobre IPv4 que usan una conexión VPN a través de Internet IPv6....	207
Figura 128: Paquetes IPv6 nativo que usan una conexión VPN a través de Internet IPv6	207
Figura 129: Método de trabajo de IPSec en IPv6, redes privadas virtuales	208
Figura 130: Método de trabajo de IPSec en IPv6, road warrior	209
Figura 131: Método de trabajo de IPSec en IPV6, túneles anidados	209
Figura 132: Mensaje DNS.	213
Figura 133: Ejemplo de petición iterada.	215
Figura 134: Ejemplo de petición recursiva.....	216
Figura 135: Funcionamiento del protocolo SSH.....	220
Figura 136: capas del estándar 802.11	222
Figura 137: Distintos estándares 802.11	223
Figura 138: Pila de protocolo SIP	226
Figura 139: Fuentes de retardo de VoIP.	227
Figura 140: Arquitectura propuesta VoIPv6.....	230
Figura 141: Tiempos de llamada vs. Paquetes IPv4 e IPv6.	231



ÍNDICE DE TABLAS

Tabla 1: Problemática IPv4.....	21
Tabla 2: Solución a IPv4.....	22
Tabla 3: Resumen de la Cabecera de Extensión IPv6.	36
Tabla 4: Compresión de la notación IPv6.....	45
Tabla 5: Formato del prefijo de una dirección Multicast.	54
Tabla 6: Formato de mensajes ICMPv6	61
Tabla 7: Protocolos IGP.....	74
Tabla 8: IGP vs EGP.	75
Tabla 9: Formato de mensaje cliente-servidor.	81
Tabla 10: Formato de DUID-LL	81
Tabla 11: Mensajes de DHCPv6.	84
Tabla 12: Formato de mensaje entre agentes de retransmisión DHCPv6.	89
Tabla 13: RIPng vs RIPv2	99
Tabla 14: Ejemplo, tabla de direccionamiento con RIPng	101
Tabla 15: Características de .EIGRPv4 vs EIGRPv6.....	105
Tabla 16: Ejemplo, tabla de direccionamiento con EIGRPv6	106
Tabla 17: Tipos de paquetes OSPF.	110
Tabla 18: Encabezado OSPFv3 vs OSPFv2.....	111
Tabla 19: Direcciones Multicast con OSPF	112
Tabla 20: LSAs (tipo9).....	114
Tabla 21: Tabla de direcciones con OSPFv6.	115
Tabla 22: Traducción de cabeceras IPv4 a cabeceras IPv6.....	164
Tabla 23: Campos de la cabecera IPv6.....	164
Tabla 24: Campos de la cabecera de fragmentación IPv6.....	165
Tabla 25: Traducción de cabeceras IPv6 a cabeceras IPv4.....	166
Tabla 26: Traducción de cabeceras IPv6 a cabeceras IPv4 (2)	167
Tabla 27: Traducción NAT-PT	171
Tabla 28: Diferencia de NAT64 con estado y sin estado.....	177
Tabla 29: Cabeceras de extensión en MIPv6.....	189
Tabla 30: Beneficios VoIPv6 respecto de VoIPv4	229
Tabla 31: Medición del retardo.	231





1. Introducción

En la actualidad el internet es el sistema global de la información, el cual se basa en el protocolo de Internet, la versión IPv4, es la primera en ser implementada en gran escala desde sus inicios hasta el día de hoy. Técnicamente este sistema de direccionamientos ya no es suficiente debido al aumento exponencial de equipos conectados a la red. Los usuarios del servicio de internet exigen nuevos alcances que el protocolo IPv4 no puede proporcionar, es limitado, y de ahí surge la necesidad de desarrollar un nuevo protocolo, el cual es denominado como protocolo IPv6.

El presente documento tiene como finalidad el estudio e implementación de IPv6 en el área de ingeniería de Telemática de la UNAN-León, brindando la oportunidad de desarrollar esquemas teóricos-prácticos, que le permitan al estudiante analizar, interpretar y adquirir conocimientos plenos de cómo funciona y actúa este protocolo. Por lo cual se realizarán prácticas guiadas de Switching y Routing con IPv6.



2. Antecedentes

En la actualidad en el Departamento de Computación de la Universidad Nacional Autónoma de Nicaragua, UNAN-León se han elaborado temas monográficos relacionados a escenario de redes, implementado IPv4.

El primer tema monográfico fue titulado **“Prácticas de Laboratorio para la Asignatura de Redes de Ordenadores II”** elaborado por la Br. Alicia Esmeralda Larios Acuña, Br. Irayda Rosa Mayorga Castellón y Br. Bruna Mercedes Moreira Cárcamo, en Mayo del 2008. El documento fue desarrollado con un total de 17 prácticas donde se abarcaron las configuraciones de switches y routers utilizando tecnología cisco, por lo que se requirió implementar el simulador Packet Tracer.

El segundo tema monográfico fue titulado **“Análisis de Software de Open Source Routing, para su uso en los laboratorios de Redes de Computadoras.”** elaborado por el Br. Léster Ramón Acevedo Hernández. Y el Br. José Miguel Bárcenas Flores, en Febrero del 2012. El documento consta de 5 prácticas donde se comprobó que la herramienta del **Open Source Routing** es la más idónea para la implementación de protocolos de enrutamientos, teniendo en cuenta los servicios IP (SSH, DHCP, NAT), Enrutamiento básico (RIP, BGP, OSPF), Servicios de seguridad (SNORT, VPN) entre otros, con el fin de ofrecer seguridad integra y funcional en el enrutamiento para entornos físicos y virtuales.

El tercer tema monográfico fue titulado **“Propuesta de prácticas de laboratorios de Switches y Routers para la carrera de Ingeniería en Telemática UNAN-León**, elaborado por el Br. Rudy Otoniel Quiróz Vázquez, el Br. Franklin Ernesto Ramírez Medina y el Br. Yoel Francisco Rivera González, en Septiembre del 2013. Este tema propone la elaboración de 15 prácticas con la finalidad de abordar temáticas teóricas y prácticas, permitiendo que los estudiantes puedan adquirir estos conocimientos, para ser puestos en práctica al desarrollar estos laboratorios.



3. Planteamiento del Problema

Según lo expuesto, es necesario que nuestro Departamento de Computación cuente con una asignatura en donde se desarrollen contenidos teóricos y prácticos del protocolo IPv6. La ausencia de un documento formal de IPv6 en nuestro Departamento, en el que se muestre una secuencia lógica de contenidos teóricos y prácticos, dificultaría la comprensión y la correcta realización de prácticas de laboratorios en la asignatura “Despliegue de IPv6”, por parte de los estudiantes.

Los contenidos desarrollados en este documento, se basan en la microprogramación de la asignatura “Despliegue de IPv6”, desarrollada en el año 2011, la cual posee las siguientes competencias:

- Comprende el direccionamiento y la asignación de en las redes de área local tomando en cuenta los equipos y las necesidades de conectividad de los usuarios.
- Comprende los mecanismos de transición y los protocolos de routing en IPv6, para su correcta implementación en redes medianas y grandes, de la red en cuestión.
- Configura y gestiona los servicios de red en IPv6, que permitan el aprovechamiento de la infraestructura de red de la empresa, tomando como referencia las necesidades de la organización

Debido a que el protocolo está comenzado a ser usado para desplegar masivamente redes LAN, WAN, redes de acceso, etc, hemos visto la necesidad de desarrollar este documento.

Todo lo planteado anteriormente acerca de IPv6 hace que surjan las siguientes preguntas generales y específicas:

Pregunta General:

- ¿Es viable que el desarrollo de contenidos teóricos y prácticos de IPv6 sea de utilidad para los estudiantes de la carrera de Ingeniería en Telemática de la UNAN-León?

Preguntas específicas:

- ¿Será necesario desarrollar un plan de prácticas de laboratorio relacionadas con Switching, Routing y Servicios de Red con IPv6 para la carrera de Ingeniería en Telemática del Departamento de Computación de la UNAN-León?
- ¿Qué secuencia debe tener el documento de manera que permita a los estudiantes poner en prácticas los conocimientos adquiridos en el transcurso del desarrollo de los laboratorios relacionados con Switching y Routing con IPv6?
- ¿Qué temas deben ser abordados, donde estos a su vez, sean de importancia en el área de redes?



4. Justificación

Teniendo como punto de referencia el agotamiento y las debilidades que presenta actualmente el protocolo IPv4, nace la idea de crear un documento, en el cual se plasmaran conocimientos teóricos y se desarrollaran prácticas de laboratorios referentes a switching y routing con IPv6, de manera eficiente, ordenada y secuencial, lo que le permitirá a los estudiantes de la carrera de Ingeniería en Telemática de la UNAN-León tener un aprendizaje fiable sobre este protocolo.

En el documento se plasmarán los siguientes aspectos: diseñar y aplicar prácticas de laboratorios con su debida documentación, enumerarlas de manera que los estudiantes vayan mejorando sus conocimientos según la complejidad de los laboratorios que se desarrollen, obteniendo mejor resultado a lo hora de aplicar su conocimientos en el ámbito profesional.

4.1 Originalidad.

Los documentos que se han elaborado son relacionados con el protocolo IPv4, con el fin de mejorar la calidad del proceso de **“enseñanza y aprendizaje”**.

El presente estudio pretende incorporar la secuencia de avanzada de IPv4, que sería IPv6, con el objetivo de aportar al Departamento de Computación de la UNAN-León una guía propia y completa, que no se dispone hasta la fecha actual, sobre el abordaje y manejo de IPv6, de cara a la nuevas exigencias de la computación moderna y del dinamismo de las empresas que requieren de nuestros servicios profesionales.

Con el protocolo IPv6 se puede crear múltiples escenarios de redes, por lo cual se desarrollarán prácticas en las siguientes áreas:

- Configuración entre IPv4 e IPv6.
- Configuración de Vlans estáticas y dinámica con IPv6.
- Configuración Frame Relay.
- Configuración Intervlans.
- Configuración de listas de acceso.
- Configuración con DNS, SSH, FTP.

4.2 Alcance.

- **Para los Docentes:** Tendrá un documento que les permita asignar prácticas guiadas a los estudiantes en forma ordenada, secuencial, según el nivel de complejidad.
- **Para los estudiantes:** Comprender los temas y solucionar las prácticas propuestas en el documento.



4.3 Producto.

El presente documento constará de 3 partes principales que se describen a continuación:

- **Guiado:** Aquí se pretende describir paso a paso la forma que se deberá realizar las prácticas, se tratará de explicar de forma fácil y sencilla el funcionamiento de los protocolos y tecnologías que se van a configurar en el desarrollo de las prácticas.
- **Sencillo:** De acuerdo al desarrollo de las prácticas y la documentación, serán elaboradas de forma que se comprenda fácilmente.
- **Secuencial:** Irá de acuerdo a la creciente dificultad y complejidad de las configuraciones de las prácticas, que se proponen de forma secuencial para un aprendizaje lógico.

4.4 Impacto.

El documento permitirá contar con un material de apoyo altamente eficiente tanto para el personal docente, como recursos en formación (estudiantes) sobre el manejo de IPv6.



5. Objetivos

5.1 Objetivo general

- Crear propuestas de prácticas de laboratorios de Switching, Routing y Servicios de Red, usando el protocolo IPv6, con la finalidad de usarlo en la asignatura “Despliegue de IPv6” correspondiente a la Electiva X, de la carrera de Ingeniería en Telemática del Departamento de Computación de la UNAN-León.

5.2 Objetivos específicos.

- Presentar un documento en el que los estudiantes adquieran conocimientos teóricos-prácticos, que les sirvan de base para el desarrollo de las prácticas de laboratorios.
- Definir el formato que regirán el enunciado de las prácticas de laboratorios a desarrollar, en base a la experiencia de formatos usados en diferentes asignaturas de la carrera.
- Definir la secuencia lógica de las prácticas de laboratorio, según la complejidad que presenta cada una ellas, abordando temas de nivel básico, medio y avanzado.

6. Diseño metodológico

La realización de este tema monográfico se realizó en diversas etapas, en los cuales se mostró los pasos para la concepción de este estudio.

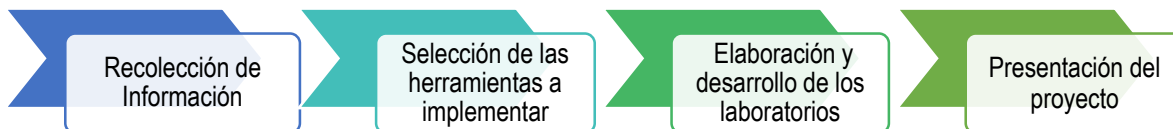


Figura 1: Etapas del trabajo.

6.1 Recolección de Información

En la primera etapa de la investigación, se realizó un estudio exhaustivo del protocolo IPv6, con la finalidad determinar los aspectos más importante a desarrollar en el tema monográfico, organizando la información según el nivel de complejidad que tiene cada uno de los temas a desarrollar en los aspectos teóricos. La secuencia de los contenidos teóricos es la siguiente:

- Definición y características del protocolo IPv6.
- Proceso de subneting en IPv6.
- Coexistencia con el protocolo IPv4.
- Implementación de IPv6 con los protocolo de enrutamiento dinámico interno y externo.
- Implementación de IPv6 con DHCP, SSH, DNS, Wireless.
- Implementación de IPv6 en la seguridad de redes.
- Entre otros temas.

6.2 Selección de las herramientas a implementar

En esta etapa se seleccionaron las plataformas a usar en el desarrollo de las prácticas de redes basadas en IPv6, obteniendo con mejor resultado el simulador Packet Tracer (Windows y Linux), GNS3 y Core Emu (Ubuntu), tomando en cuenta el soporte que tiene cada uno de ellos para las tecnologías y/o protocolos que serán usados .



Figura 2: Simuladores usados.

6.3 Elaboración y desarrollo de los laboratorios

- Organización de las prácticas: Es el punto donde la información es organizada según el nivel de complejidad que tienen los temas a desarrollar en los aspectos prácticos. El orden de las prácticas a desarrollar es el siguiente:



- Direccionamiento IPv6 con rutas estáticas.
 - DHCPv6 y Autoconfiguración.
 - Coexistencia de IPv4 e IPv6.
 - VLANs estáticas y dinámicas.
 - Frame Relay e Intervlans.
 - Listas de acceso.
 - Enrutamiento interno: RIPng, OSPFv3, EIGRPv6, IS-ISv6.
 - Enrutamiento dinámico interno y externo: RIPng, OSPFv3, EIGRPv6 + BGP4.
 - VPN IPv6
 - VoIP IPv6
 - DNS.
 - HTTP.
 - SSH.
 - Miscelánea.
- Desarrollo del enunciado de prácticas: El formato a seguir para enunciar cada una de las prácticas propuestas es el siguiente:

Titulo

Nombre de la práctica.

Objetivos

- Presenta una visión general de lo que se espera lograr con el desarrollo de la práctica.
- Expondrá aspectos específicos, en los cuales los estudiantes deberán de enfocar su trabajo de laboratorio.

Introducción

Contiene rasgos generales de lo que posee cada práctica en el desarrollo de su contenido, y en algunos casos aspectos claves que los estudiantes deben tomar en cuenta para facilitar la solución de las mismas.

Requerimientos

- **Hardware:** Contiene los requerimientos que se deben tomar en cuenta para la realización de las prácticas.
- **Software:** Define el entorno en que se desarrollara la práctica.

Conocimientos previos

Serán detallados los conocimientos mínimos que deberá tener el estudiante para poder dar solución a la práctica enunciada. En algunos casos se hará referencia a prácticas antes enunciadas y documentación extra de ser necesario.



Topología

Se detallara mediante una imagen donde se represente la topología correspondiente a la práctica.

Funcionalidad

Explica de manera general la funcionalidad de la figura que es mostrada como topología.

Resumen de comandos

Presenta una lista de comandos, los cuales serán de ayuda para dar una solución correcta a la práctica.

Datos de los equipos

Los datos de los equipos serán mostrados en tablas o cuadros.

Enunciado

Expone de forma clara las configuraciones que se deberán hacer en cada equipo para poder dar una correcta solución.

Tiempo estimado de solución

Tiempo estimado en horas presenciales y no presenciales para dar solución a cada práctica, con la salvedad que una vez que sea usada por diferentes maestros puede variar de acuerdo al criterio personal de evaluación de cada maestro.

Preguntas de análisis

Se evalúa grado de asimilación y comprensión de los conceptos básicos y configuraciones realizadas después de resolver cada práctica.



CRONOGRAMA DE ACTIVIDADES																																				
Actividades	2014												2015																							
	Octubre				Noviembre				Diciembre				Enero				Febrero				Marzo				Abril				Mayo				Junio			
	1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4	1	2	3	4
Selección del Tema	█																																			
Introducción		█																																		
Antecedentes		█																																		
Definición de Problema		█																																		
Justificación			█																																	
Objetivos			█																																	
Diseño Metodológico			█																																	
Investigación				█	█	█	█	█	█	█	█																									
Marco Teórico												█	█	█	█	█	█	█	█	█																
Desarrollo Práctico																				█	█	█	█	█	█	█	█	█	█							
Conclusión																																	█			
Recomendaciones																																	█			
Bibliografía																																	█			
Presentación del proyecto																																			█	





1. IPv4

1.1 Definición

El protocolo IPv4 es la cuarta versión del protocolo de internet, se basa en la transmisión de datos entre dispositivos a través de la conexión de redes, este protocolo es el más usado en el modelo TCP/IP. Los datos viajan a través de innumerables redes físicas, para lograr que la información sea enviada y recibida correctamente, los datos son encapsulados en lo que denominamos paquete de datos o datagrama IP, los cuales incluye campos o segmentos que permite controlar el funcionamiento del protocolo de Internet, y asegurar que los datos lleguen a su destino.

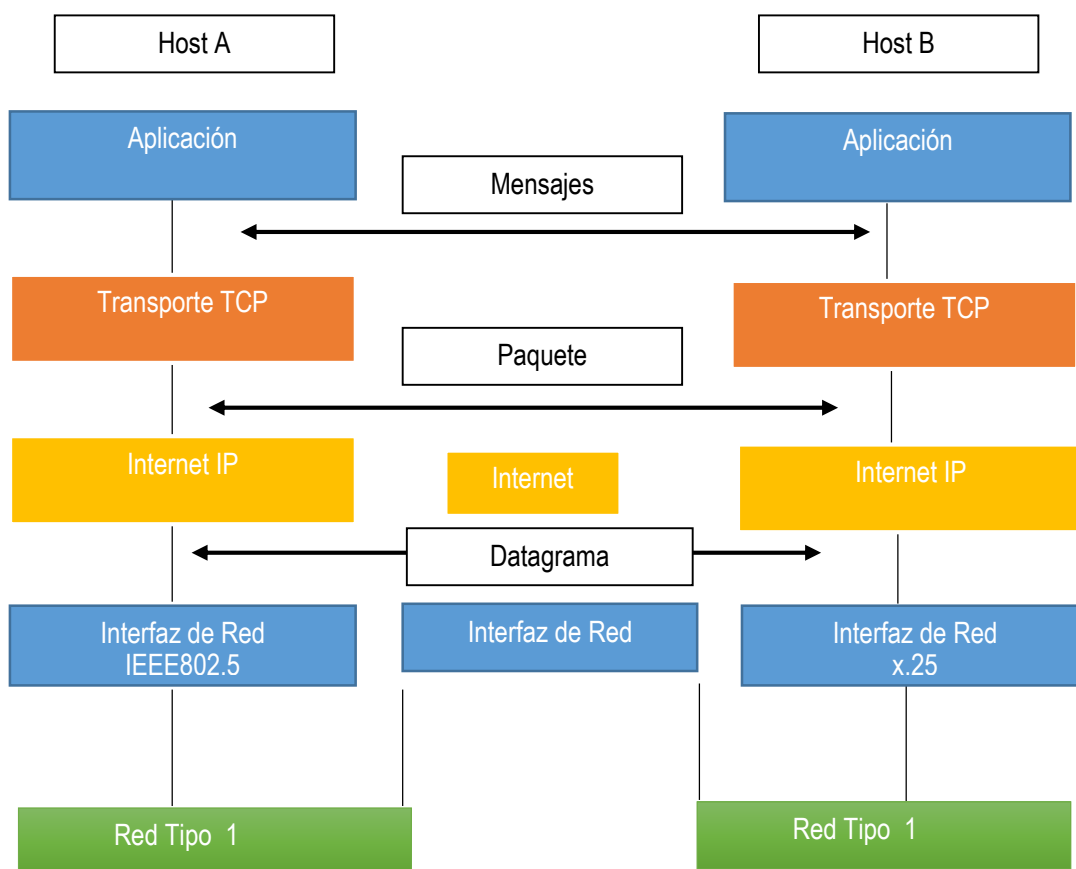


Figura 3: Comunicación con diferentes sistemas de internet.



1.2 Paquete o Datagrama

La estructura de un datagrama IPv4, se divide en un bloque de 32 bits, conceptualmente se divide en dos partes: Cabecera IP y campo de Dato.

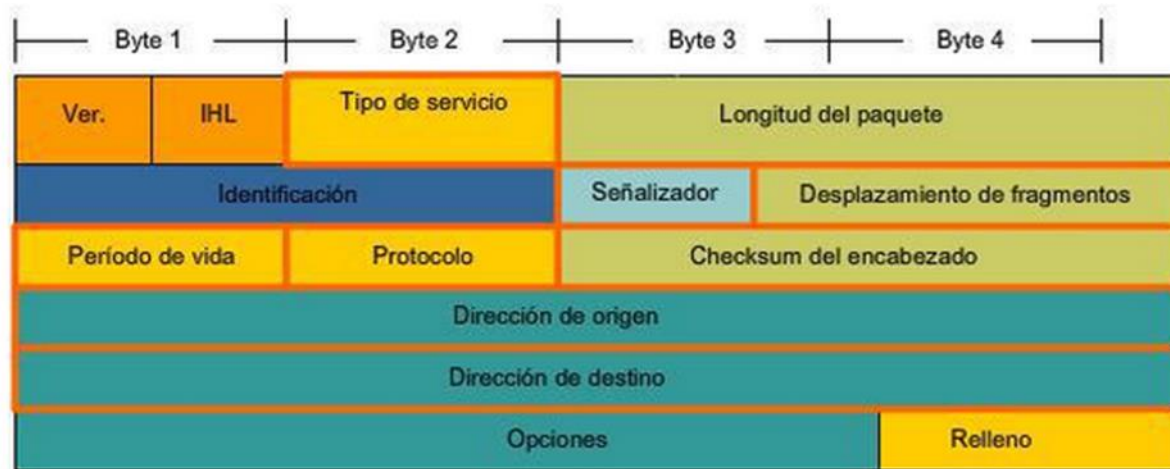


Figura 4: Datagrama IPv4.

- **Versión:** Especifica la versión del protocolo IP a la que pertenece el datagrama, actualmente se está utilizando la versión 4 del protocolo.
- **Longitud de la cabecera:** es un campo de 4 bits, que proporciona la longitud del encabezado del datagrama, medido en palabras de 32 bits, su valor mínimo es de 5 palabras ($5 \times 32 = 160$ bits, 20 bytes) para una cabecera correcta, y el máximo de 15 palabras ($15 \times 32 = 480$ bits, 60 bytes).
- **Tipo de servicio:** especifica prioridad y tipo de transporte.
- **Longitud Total:** proporciona la longitud del datagrama medido en bytes, incluyendo los bytes del encabezado y los datos.
- **Identificación:** Es un entero de 16 bits que identifica al datagrama y lo distingue de otros datagramas que hemos enviado. Es una especie de número de secuencia que se incrementa cada vez que IP envía un datagrama.
- **Flags + Desplazamiento de fragmento:** Estos campos incluyen información útil para el mecanismo de fragmentación de datagramas. Cuando un datagrama cruza una pasarela y al otro lado existe una red con un MTU inferior al tamaño del datagrama, la pasarela lo fragmenta en trozos. Estos fragmentos son datagramas que viajan hacia el destino de forma independiente, donde son recogidos por el protocolo IP para reconstruir el datagrama original.
- **Tiempo de vida:** Especifica la duración en segundos del tiempo que el datagrama tiene permitido permanecer en la red.



- **Protocolo:** Especifica que protocolo de alto nivel se empleó para construir el mensaje transportado en el campo datos de Datagrama IP. Algunos valores posibles son: 1= ICMP, 6 = TCP, 17 = UDP, 88 = IGRP (Protocolo de Enrutamiento de Pasarela Interior de CISCO).
- **Checksum de la cabecera:** En este campo se almacena un Checksum de los campos de la cabecera. Es un mecanismo simple para detectar posibles errores en los campos de la cabecera del datagrama, los cuales podrían provocar situaciones "incómodas" en la red.
- **Direcciones IP origen y destino:** Direcciones origen y destino del datagrama. Aunque el datagrama viaje a través de varias pasarelas, estos campos nunca cambian.
- **Opciones IP:** Este campo es opcional y de longitud variable, se incluye en principio para pruebas de red o depuración.
- **Relleno:** Campo utilizado para completar el tamaño de 32 bits en el área de datos.

1.3 Direccionamiento.

Las direcciones IP son números binarios de 32 bits que son usados como direcciones en los protocolos IPv4. Las direcciones IP se escriben mediante la denominada notación punto decimal o de cuatro octetos, cada octeto está formado por bytes, entre 0 y 255, y tiene el siguiente formato XXXX.XXXX.XXXX.XXXX. Por lo cual los equipos en una red usan estas direcciones para comunicarse, de manera que cada equipo tenga una IP exclusiva. Las direcciones IP proporcionan dos datos: el número de red y el número de host. Para que un sistema pueda transmitir datos, debe determinar con claridad la dirección destino de red y host, además de poder informar al resto de sistemas de cuál es su propia dirección de red y host.

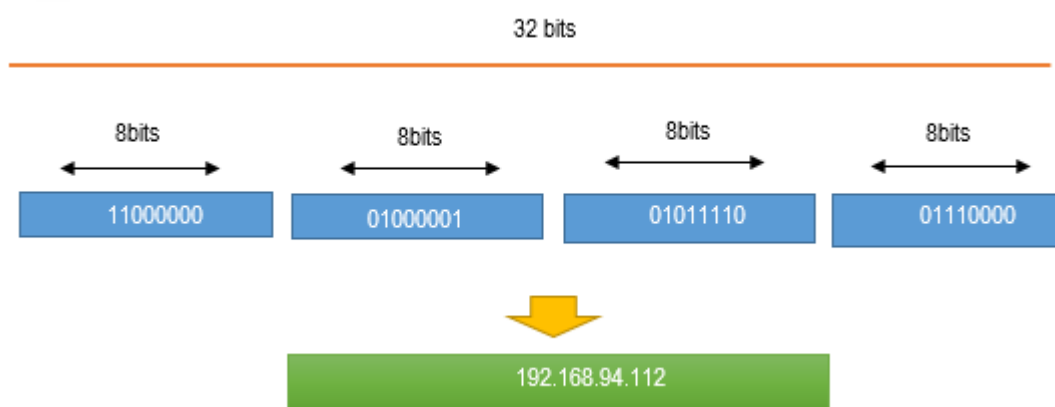


Figura 5: Ejemplo de dirección IPv4.

Los sistemas de red se pueden direccionar de tres formas:

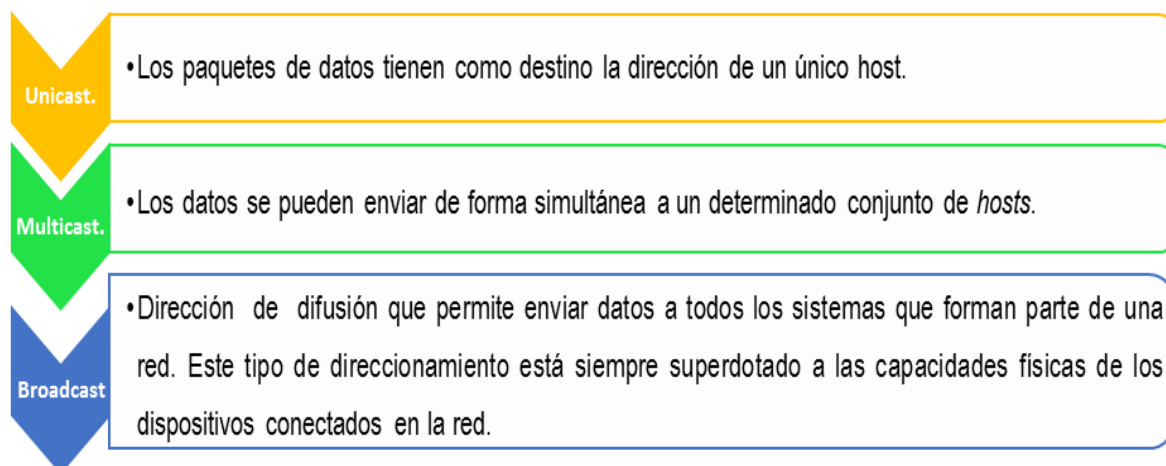


Figura 6: Formas de direccionar una dirección IPv4

1.4 Clases de direccionamiento.

El protocolo IPv4 se basa en la arquitectura de clases, lo que facilita tener diferentes tipos de direcciones IP dependiendo del tamaño de la red, existen cuatro formatos para la dirección IPv4: Clase A, B, C y D.

Clase A: Si el bit de mayor peso es «0» la máscara por defecto tendrá un prefijo de 8 bits. Se tienen por tanto 8 bits para direcciones de red y 24 bits hosts. Las direcciones de clase A están concebidas para redes compuestas por numerosos ordenadores. Puesto que son escasas las redes de estas características, se dedican pocos bits para identificar la red; sólo siete bits que permiten numerar hasta 27, es decir 128 redes.

Clase B: Si los dos primeros bits son «1» «0», la máscara por defecto tendrá una longitud de 16 bits (prefijo 16). Con ello los primeros 16 bits son para identificar la red y los 16 últimos es para identificar los hosts. Estos tipos de direcciones se emplean en redes constituidas por número medios de ordenadores. Se produce la circunstancia de que existe un número también intermedio de estas redes (se permite hasta 214, es decir 16,384 redes de esta clase).

Clase C: Si los tres primeros bits son «1» «1» «0» la máscara por defecto tiene un prefijo de 24 bits. Para esta clase se contempla la existencia de una gran cantidad de redes, en concreto 224. En cada una de ellas el número de equipos es como máximo 253, una vez restadas las direcciones de red y difusión. Las direcciones de esta clase se destinan a redes con pocos ordenadores, que son lo más frecuentes.

Clase D: Si los cuatro primeros bits de la dirección son «1» «1» «1» «0» nos encontramos frente a una dirección **Multicast**. Entonces, no se habla de una dirección de red, sino de un grupo de equipos a los que



se desea enviar datos simultáneamente. Todos los bits de una dirección *Multicast* son significativos, así que la máscara por defecto es de 32 bits (prefijo 32).

Clase E: Si los cuatro primeros bits de la dirección son unos lógicos, la dirección IP pertenece a un rango que se ha reservado para experimentación. Dentro de esta clase aparece la dirección IP de difusión 255.255.255.255.

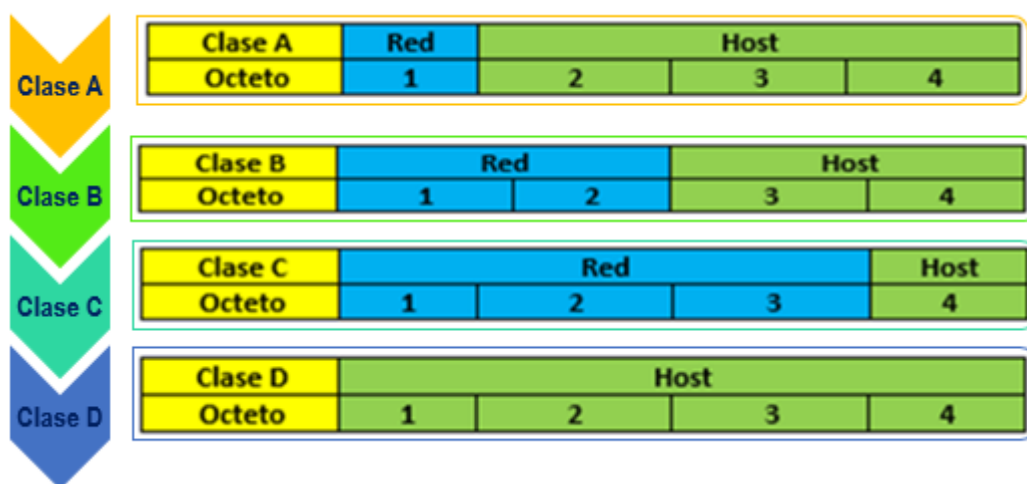


Figura 7: Distribución de bits de red y host en las diferentes clases de IPv4.

RANGOS IP		
Clase A	0.0.0.0	00000000.00000000.00000000.00000000
	126.255.255.255	01111110.11111111.11111111.11111111
Clase B	128.0.0.0	10000000.00000000.00000000.00000000
	191.255.255.255	10111111.11111111.11111111.11111111
Clase C	192.255.255.255	11000000.11111111.11111111.11111111
	223.255.255.255	11011111.11111111.11111111.11111111
Clase D	224.0.0.0	11100000.00000000.00000000.00000000 (Multicast)
	239.255.255.255	11101111.11111111.11111111.11111111
Clase E	240.0.0.0	11110000.00000000.00000000.00000000 (Experimentales)
	255.255.255.254	11111111.11111111.11111111.11111110

Figura 8: Clase de direccionamiento IPv4.



1.5 Ruteo interno de dominio sin clases (CIDR).

A medida que el internet comenzó a crecer de manera espectacular, surgieron problemas en el direccionamiento de "Clases" como son:

- La falta de flexibilidad en el direccionamiento interno.
- El mal uso del espacio de direcciones.

Estas dificultades fueron solucionadas temporalmente a través del direccionamiento de sub redes. Con el fin de extender la vida del protocolo IPv4, era necesario adoptar un nuevo enfoque, este sistema implica la eliminación de clases de direcciones por completo, creando un nuevo esquema de direccionamiento sin clases llamado Classless Inter-Domain Routing (CIDR).

Classless Inter-Domain Routing (CIDR) es un sistema de direccionamiento y enrutamiento IP, el cual representa las direcciones y máscaras de subred en notación binaria, con el fin de dividir los tamaños de red fijos tradicionales, permitiendo tener una elección más eficaz para la asignación de direcciones, esto representa una mejora en el modo de interpretar las direcciones.

CIDR ofrece numerosas ventajas frente al esquema de direccionamiento "clases", se utilicen o no subredes:

Asignación eficiente del espacio de direcciones: En lugar de asignar direcciones en bloques de tamaño fijo de granularidad baja, las direcciones CIDR son asignadas en tamaños de cualquier múltiplo binario.

Eliminación del desequilibrio de clases: No hay más redes de clases A, B y C, de manera que no hay problemas con que algunas porciones del espacio de direcciones se utilicen ampliamente.

Método de subnetting sin divisiones: CIDR implementa los conceptos de las subredes dentro de la propia Internet. Una organización puede utilizar el mismo método utilizado en Internet para subdividir la red interna en subredes de complejidad arbitraria sin necesidad de un mecanismo de división en subredes separadas.

Las direcciones IP están diseñadas para ser divididas en identificador de red e identificador de host. Entonces, cuando se introdujeron las subredes, "robamos" bits del ID de host para crear un identificador de subred, dando a la dirección IP un total de tres niveles jerárquicos. Con VLSM, subnetearnos más aun las subredes, tomando más bits del ID de host lo que resultó en una jerarquía de varios niveles con "sub-subredes", "sub-sub-subredes" y así sucesivamente.

En un entorno sin clases, cambiamos completamente la forma en que vemos las direcciones IP, mediante la aplicación de conceptos de VLSM no sólo a una red, sino a la totalidad de Internet. En esencia, el Internet se convierte en una sola red gigante que es "subnetead" en una serie de grandes bloques.

Algunos de estos grandes bloques se estructurarán en bloques más pequeños, que a su vez pueden dividirse aún más. Esta división puede ocurrir varias veces, lo que nos permite dividir las direcciones de Internet en diferentes tamaños, para adaptarlas a las necesidades de las empresas.



Notación CIDR

Así como en el subnetting requerimos del uso de una máscara de subred para mostrar que partes pertenecen a la ID de red o al ID de subred y cual al ID de host, CIDR usa una máscara de subred para mostrar dónde se traza la línea entre el ID de host y el ID de red. Sin embargo, por simplicidad, bajo CIDR no sobemos trabajar con máscaras de subred binarias de 32 bits. En su lugar, usamos la notación de barra, más propiamente llamada notación CIDR. En este método, se muestra el tamaño de la red, a veces llamada longitud del prefijo, siguiendo la dirección IP de un número entero que nos dice cuántos bits se utilizan para la identificación de la red (prefijo).

Bloque CIDR

CIDR facilita el encaminamiento al permitir agrupar bloques de direcciones en una sola entrada de la tabla de rutas. Estos grupos, llamados comúnmente Bloques CIDR, comparten una misma secuencia inicial de bits en la representación binaria de sus direcciones IP. Los bloques CIDR IPv4 se identifican usando una sintaxis similar a la de las direcciones IPv4: cuatro números decimales separados por puntos, seguidos de una barra de división y un número de 0 a 32; A.B.C.D/N.

Los primeros cuatro números decimales se interpretan como una dirección IPv4, y el número tras la barra es la longitud de prefijo, contando desde la izquierda, y representa el número de bits comunes a todas las direcciones incluidas en el bloque CIDR. Dado que la longitud de una dirección IPv4 es fija, de 32 bits, un prefijo CIDR de N-bits deja bits sin encajar, y hay combinaciones posibles con los bits restantes. Esto quiere decir que direcciones IPv4 encajan en un prefijo CIDR de N-bits. Los prefijos CIDR cortos (números cercanos a 0) permiten encajar un mayor número de direcciones IP, mientras que prefijos CIDR largos (números cercanos a 32) permiten encajar menos direcciones IP. Una dirección IP puede encajar en varios prefijos CIDR de longitudes diferentes.

CIDR también se usa con direcciones IPv6, en las que la longitud del prefijo varía desde 0 a 128, debido a la mayor longitud de bit en las direcciones, con respecto a IPv4.

CIDR y máscara de subred.

Una máscara de subred es una máscara que codifica la longitud del prefijo de una forma similar a una dirección IP - 32 bits, comenzando desde la izquierda, ponemos a 1 tantos bits como marque la longitud del prefijo, y el resto de bits a cero, separando los 32 bits en cuatro grupos de ocho bits.

CIDR usa máscaras de subred de longitud variable (VLSM) para asignar direcciones IP a subredes de acuerdo a las necesidades de cada subred. De esta forma, la división red/host puede ocurrir en cualquier bit de los 32 que componen la dirección IP. Este proceso puede ser recursivo, dividiendo una parte del espacio de direcciones en porciones cada vez menores, usando máscaras que cubren un mayor número



de bits. Las direcciones de red CIDR/VLSM se usan a lo largo y ancho de la Internet pública, y en muchas grandes redes privadas.

Agregación de prefijos.

Otro beneficio de CIDR es la posibilidad de agregar prefijos de encaminamiento, un proceso conocido como "supernetting". Por ejemplo, dieciséis redes /24 contiguas pueden ser agregadas y publicadas en los enrutadores de Internet como una sola ruta /20 (si los primeros 20 bits de sus respectivas redes coinciden). Dos redes /20 contiguas pueden ser agregadas en una /19, etc.

Esto permite una reducción significativa en el número de rutas que los enrutadores en Internet tienen que conocer (y una reducción de memoria, recursos, etc.) y previene una explosión de tablas de encaminamiento, la dirección menor (más baja - todos los bits de host a 0) del bloque se usa para identificar a la propia red (toda la red), y la dirección mayor (la más alta - todos los bits de host a 1) se usa como dirección de Broadcast. Por tanto, en un bloque CIDR /24 podríamos disponer de direcciones IP para asignar a dispositivos.



1.6 Problemas con IPv4

La necesidad de un espacio de direcciones extenso está forzando a un cambio inmediato en el Protocolo de Internet, debido a las limitaciones que presenta el protocolo actual.

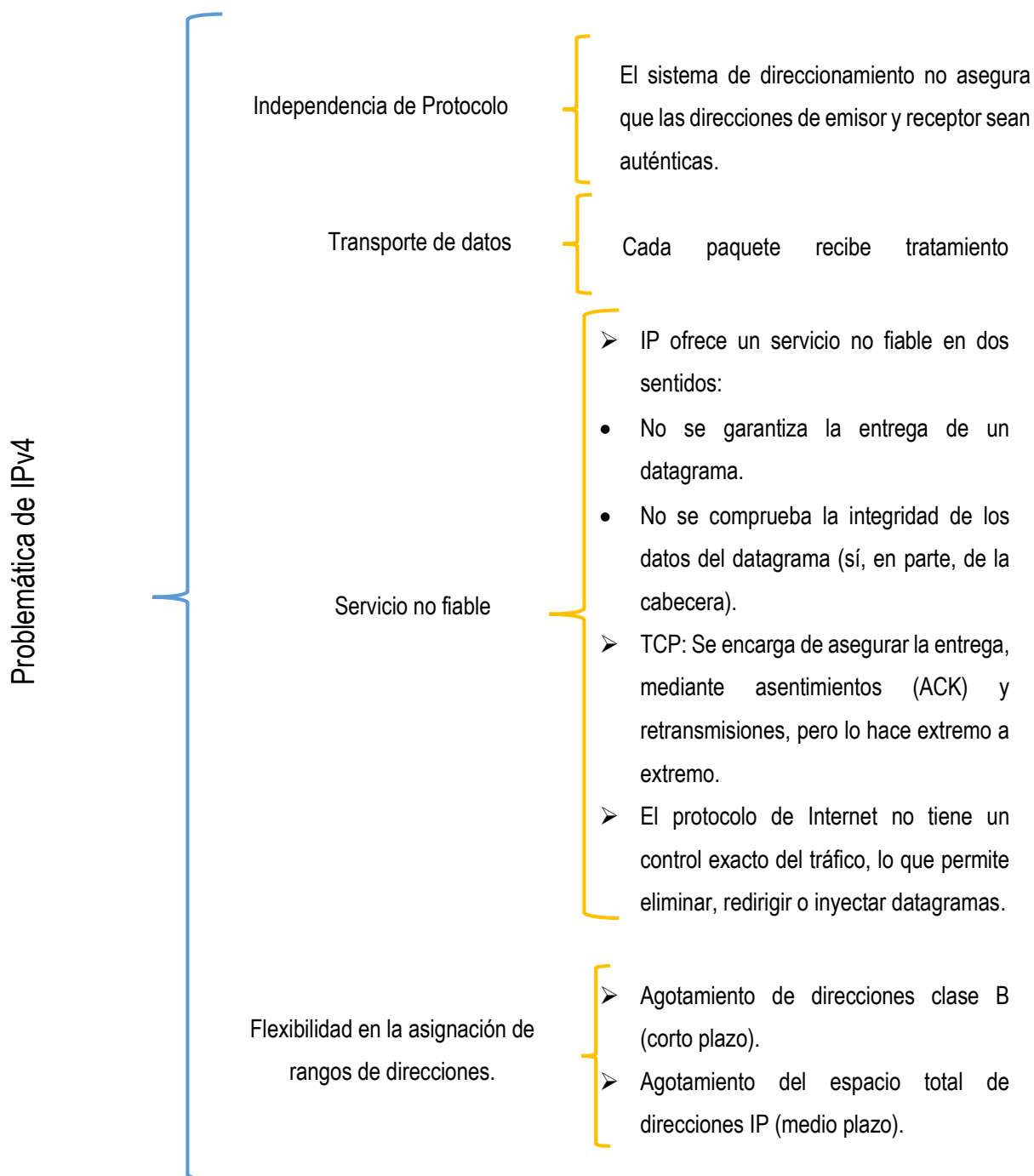


Tabla 1: Problemática IPv4.



1.7 Solución a IPv4

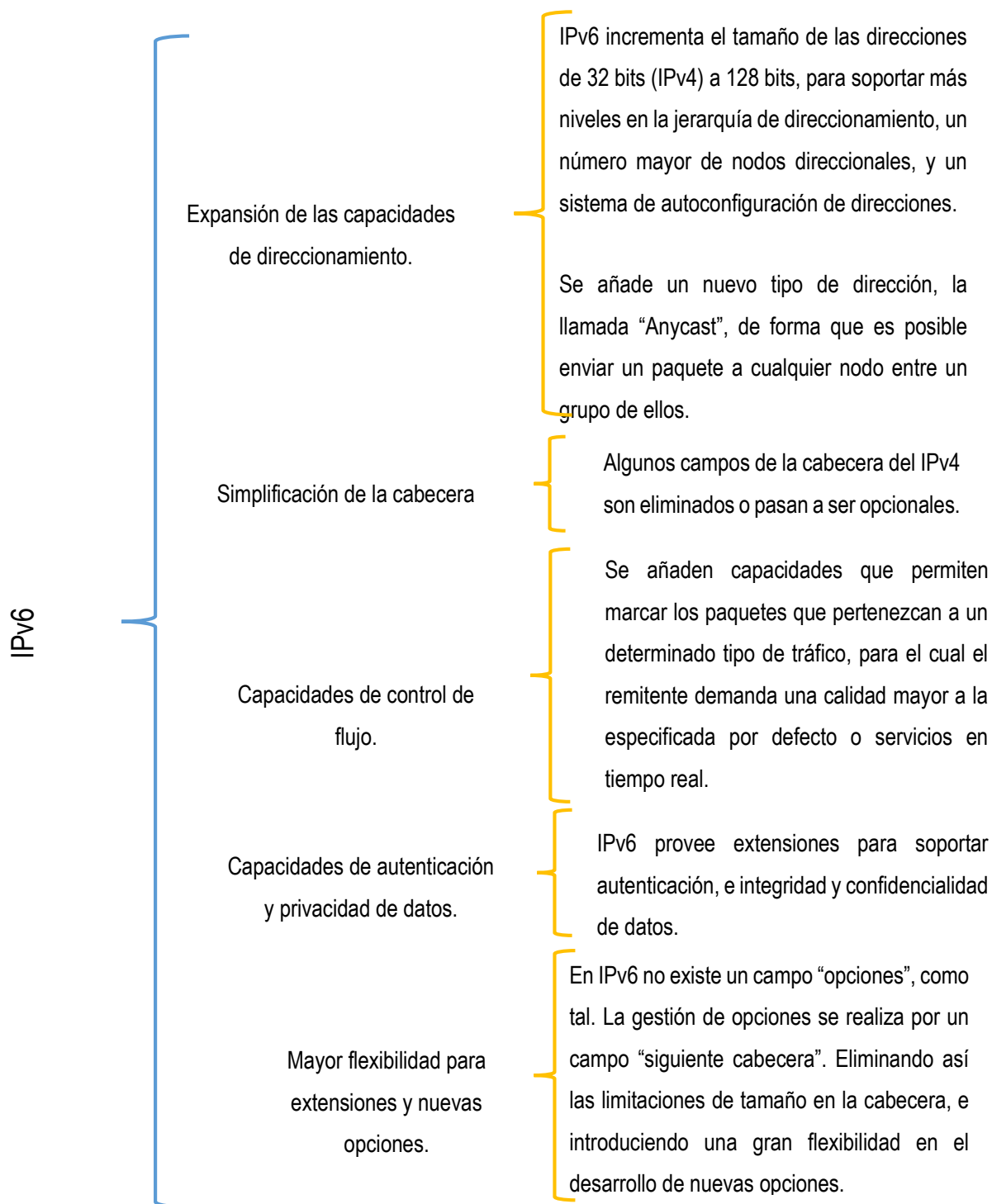


Tabla 2: Solución a IPv4.



1.8 ¿Por qué cambiar a IPv6?

La versión 4 del protocolo de Internet IP proporciona los mecanismos de comunicación básicos del conjunto TCP/IP y de la red global del Internet; ha probado ser un diseño flexible y poderoso. Desde el momento en que se diseñó el protocolo IPv4, las tecnologías LAN, WLAN y WAN han emergido y el número de anfitriones en Internet ha crecido gradualmente. La tecnología básica TCP/IP ha funcionado bien por mucho tiempo. ¿Por qué debería cambiarse?

Se han agotado las direcciones IPv4.

IPv6 ha sido diseñado para ser fácil: Una de las características de IPv6 es que sea transparente para los usuarios y especialmente en cuanto a la configuración de sus redes y dispositivos, es lo que denominamos “autoconfiguración”.

Hay que recuperar la conectividad extremo a extremo: una de las medidas adoptadas ha sido el uso de traductores de direcciones (NAT). Estos traductores de direcciones implican que no es posible la conexión directa extremo a extremo y como consecuencia, solo funcionan correctamente en las aplicaciones cliente-servidor, y por tanto Internet se ha convertido en una red mucho más compleja, cara y difícil de gestionar.

Es necesario incrementar la seguridad en Internet: A menudo se asocia NAT a seguridad, lo cual es incorrecto. La mejor seguridad la suministran los cortafuegos y otros dispositivos especializados. IPv6 no es más seguro que IPv4, sin embargo el estándar obliga a incorporar el protocolo IPsec (seguridad IP), y al no requerir NAT, se puede utilizar IPsec extremo-a-extremo, lo cual puede utilizarse para incrementar la seguridad en la Red.

Porque disponemos de un número casi ilimitado de direcciones: Con IPv4 cada usuario recibe una única dirección, que sólo la puede utilizar el encaminador (routers) o NAT. En cambio, con IPv6, cada usuario, reciben un conjunto de direcciones, con un prefijo /48, es decir, de 48 bits.

Se pueden utilizar sistemas de multidifusión: Los sistemas de multidifusión (Multicast) también son posibles con IPv4, pero son mucho más costosos y complicados de manejar. Con multidifusión IPv6 aprovecharemos mejor la capacidad de las redes para servicios de valor añadido de vídeo y audio sobre redes de banda ancha.

El protocolo IPv6 no tiene límites.



1.9 Ejercicios de IPv4.

1. ¿Cuál de los siguientes es una dirección Clase C IP válida que se puede asignar a un host?

- a) 1.1.1.1
- b) 200.1.1.1
- c) 128.128.128.128
- d) 224.1.1.1
- e) 223.223.223.255

2. ¿Cuál es el rango de los valores asignables para el primer octeto para las redes de Clase A de propiedad intelectual?

- a) 0 a 127
- b) . 0-126
- c) 1 a 127
- d) 1-126
- e) 128-191
- f) 128-192

3. PC1 y PC2 están en dos LANs Ethernet diferentes que están separadas por un routers IP. La dirección IP de la PC1 es 10.1.1.1, y no se utiliza la división en subredes. ¿Cuál de las siguientes direcciones se podrían utilizar para PC2? (Elegir dos respuestas).

- a) 10.1.1.2
- b) 10.2.2.2
- c) 10.200.200.1
- d) 9.1.1.1
- e) 225.1.1.1
- f) 1.1.1.1

4. ¿Cuál de las siguientes opciones son verdad acerca de un host TCP / IP conectado a la LAN y su enrutamiento IP (reenvío)? (Elegir dos respuestas.)

- a) El anfitrión siempre envía paquetes a su puerta de enlace predeterminada.
- b) El host envía paquetes a su puerta de enlace predeterminada si la dirección IP de destino pertenece a una clase diferente de red IP que el host.
- c) El host envía paquetes a su puerta de enlace predeterminada si la dirección IP de destino se encuentra en una subred diferente a la de acogida.
- d) El host envía paquetes a su puerta de enlace predeterminada si la dirección IP de destino está en la misma subred que el host.



5. **¿Cuál de las siguientes son las funciones de un protocolo de enrutamiento? (Elegir dos respuestas.)**
- a) Dar a conocer las rutas a los routers vecinos
 - b) Rutas de aprendizaje para subredes directamente conectadas al routers
 - c) Aprender rutas, y poner esas rutas en la tabla de enrutamiento, para las rutas anunciadas al routers por sus routers vecinos.
 - d) Reenvío de paquetes IP basados en la dirección IP de destino de un paquete.
6. **Una empresa implementa una red TCP / IP, con un equipo PC1 en una LAN Ethernet. ¿Cuál de los siguientes protocolos y características requiere la PC1 para aprender la información de algún otro dispositivo del servidor?**
- a) ARP
 - b) Ping
 - c) DNS
 - d) Ninguna de las otras respuestas son correctas
7. **¿Cuál de las siguientes son las funciones de la capa 3 del sistema OSI? (Elegir dos respuestas.)**
- a) Direccionamiento lógico.
 - b) Direccionamiento físico.
 - c) Selección de trazado.
 - d) Arbitraje.
 - e) Error recovery



2. IPv6

2.1 Introducción

El protocolo IPv6 tiene la habilidad de escalar redes para futuras demandas que requieren fuentes ilimitadas de direcciones IP, el protocolo de internet versión 6 combina direccionamiento extendido con un encabezado más eficiente y de mejores características que IPv4. Actualmente IPv6 se conoce por “IP Next Generation” o “IPng”.

IPv4 había resultado ser un protocolo completo y de fácil implementación. El problema es que no se anticiparon algunas situaciones que eventualmente se convertirían en limitantes para la utilización del mismo:

- El crecimiento desmedido del Internet y la reducción del espacio para asignar direcciones IP.
- La necesidad de una configuración simple.
- La necesidad de una mayor seguridad a nivel IP.
- La necesidad de un mejor soporte en la transmisión de datos en “tiempo real”, mejor conocido como “Calidad de Servicio”.

Por lo antes mencionado el protocolo no presenta un futuro alentador para las redes que están por desarrollarse, por lo que se pretende que IPv6 cumplan con las exigencias de innumerables requerimientos de direccionamientos jerárquicos que el protocolo de Internet versión 4 no proporciona, este nuevo protocolo permite la comunicación de extremo a extremo sin la necesidad de la traducción de direcciones de red (Network Address Translation -NAT), lo que permite tener una nueva generación de experiencias compartidas y aplicaciones en tiempo real.

El internet se transformara a IPv6 reemplazando IPv4, sin embargo esta transformación será gradual, por lo cual IPv4 no desaparecerá de la noche a la mañana, más bien lo que se pretende es que IPv4 pueda coexistir con IPv6 mientras se logra una transición total a este nuevo protocolo. Los cambios en IPv6 se realizaron principalmente en dos aspectos: Ampliación del campo de dirección IP a 128 bits, aumentando de 32 a 128 bits cada dirección y el campo de longitud fija, para facilitar el proceso que se da a cada datagrama en los routers para encaminarlo hacia su destino. En esencia el protocolo IPv6 sigue teniendo las mismas características de la versión 4, un protocolo no fiable y no orientado a la conexión, el servicio que presta funciona y es lo suficientemente flexible para las necesidades de hoy en día y delega la confiabilidad a los protocolos superiores que permiten mantener las capas del modelo TCP/IP.



2.2 Características.

IPv6 presenta ciertas características que contrastan con la versión 4 de este protocolo. Estas características se listan a continuación:



Figura 9: Características de IPv6

Mayor espacio de direccionamiento

IPv6 utiliza direcciones de origen y destino de 128 bits (16 bytes). Aunque con 128 bits se pueden proporcionar más de $3,4 \times 10^{38}$ combinaciones posibles, el amplio espacio de direcciones de IPv6 se ha diseñado para permitir múltiples niveles de división en subredes, asignación de direcciones de la red troncal Internet a las subredes individuales de una organización. Aunque actualmente sólo un pequeño porcentaje de direcciones posibles se asignan para el uso de hosts, hay disponibles muchas direcciones para su uso en el futuro. Al tener un número mucho mayor de direcciones disponibles, ya no son necesarias las técnicas de conservación de direcciones, como la implementación de NAT.

Simplificación de cabecera

La cabecera IPv6 fue modificada para disminuir el tiempo que tardaban los enrutadores en procesarla. Esto se logró eliminando algunos campos obsoletos y moviendo los campos opcionales y los que no se consideraban indispensables a las cabeceras de extensión, las cuales se colocan después de la cabecera IPv6.



Cabecera de Extensión

IPv6 puede ser expandido para soportar nuevas características, agregando cabeceras de extensión después de la cabecera IPv6. A diferencia del campo de “Opciones” de la cabecera IPv4, el cual solo puede contener 40 bytes, el tamaño de las cabeceras de extensión es limitado únicamente por el tamaño del paquete IPv6.

Mejora de la Compatibilidad para la calidad de servicios (QoS)

Los nuevos campos del encabezado IPv6 definen cómo se controla e identifica el tráfico. La identificación del tráfico, mediante un campo Flow Label (etiqueta de flujo) en el encabezado, permite que los enrutadores identifiquen y proporcionen un control especial de los paquetes que pertenecen a un flujo dado. Un flujo es un grupo de paquetes entre un origen y un destino. Dado que el tráfico está identificado en el encabezado IPv6, la compatibilidad con QoS se puede obtener de forma sencilla incluso si la carga del paquete está cifrada con IPsec.

Mayor seguridad en el protocolo

La compatibilidad con IPsec es un requisito del conjunto de protocolos IPv6. Este requisito proporciona una solución basada en estándares para las necesidades de seguridad de red y aumenta la interoperabilidad entre diferentes implementaciones de IPv6.

Direccionamiento jerárquico y enrutamiento eficientes

Las direcciones IPv6 globales utilizadas en la porción IPv6 del Internet fueron diseñadas para crear una infraestructura de enrutamiento eficiente y jerárquica, basada en la existencia de diferentes proveedores de servicio de Internet, cada uno con diferentes características. Debido a estas características, en la parte IPv6 del Internet los enrutadores pertenecientes al backbone manejan tablas de enrutamiento mucho más pequeñas.

Nuevo protocolo para la interacción de nodos vecinos

El protocolo descubrimiento de vecinos en IPv6 consiste en un conjunto de mensajes del Protocolo de mensajes de control de Internet para IPv6 (ICMPv6, Internet Control Message Protocol for IPv6) que administran la interacción de nodos vecinos (es decir, nodos que se encuentran en el mismo vínculo). El descubrimiento de vecinos reemplaza los mensajes de Protocolo de resolución de direcciones (ARP, Address Resolution Protocol), Descubrimiento de enrutadores ICMPv4 y Redirección ICMPv4 con mensajes eficaces de multidifusión y unidifusión, y proporciona funciones adicionales. IPv6 se puede ampliar con nuevas características al agregar encabezados de extensión a continuación del encabezado IPv6. A diferencia del encabezado IPv4, que sólo admite 40 bytes de opciones, el tamaño de los encabezados de extensión IPv6 sólo está limitado por el tamaño del paquete IPv6.

2.3 Datagrama IPv6

Forma general.

La unidad de datos del protocolo IPv6 también llamada

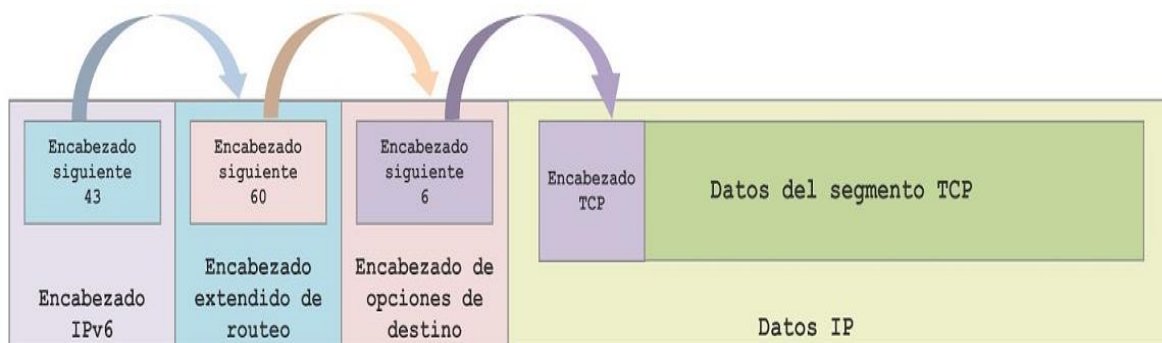


Figura 10: Forma general de un datagrama de IPv6.

IPv6 utiliza una cabecera principal con la información esencial para el encaminamiento de los datagramas, el resto de la información se incluye de forma opcional en cabeceras secundarias, las cuales se pueden o no enviar. Esta característica facilita el procesamiento de los datagramas en los routers.

Formato del encabezado base de IPv6

El protocolo IPv6 cambia completamente el formato del datagrama, uno de los elementos a considerar, es que el encabezado IPv4 tiene una longitud de 20 bytes y el encabezado IPv6 tiene una longitud de 40 bytes y la estructura de este nuevo protocolo se ha simplificado (retirando o agregando nuevos campos). Como podemos observar en la siguiente imagen:

Versión	Clase (Class)	Tipo de Flujo	
Tamaño de los datos (Payload Length)	Siguiente Cabecera (Next Header)	Alcance del Datagrama (Hop Limit)	
Dirección de Origen 128 bits (Source Address)			
Dirección de Destino 128 bits (Source Address)			
Datos			

Figura 11: Formato de Encabezado de IPv6.

- a) **Versión:** Es un campo de 4 bits que permanece al principio del encabezado e indica la versión del protocolo, que este caso es la versión 6.



- b) **Traffic Class:** Campo de 8 bits y sus funciones son similares al de tipo de servicio en IPv4. Este campo etiqueta el paquete IPv6 con un Punto de Código de Servicios Diferenciados (DSCP) que especifica cómo debe ser manejado el paquete.
- c) **Tipo de Flujo:** Un campo de 20 bits de longitud completamente nuevo que permite marcar flujos de tráfico correspondientes a diferentes conversaciones con un valor único, lo que permite darle al tráfico un tratamiento por flujos sin necesidad de revisar los encabezados correspondientes a capas superiores.
- d) **Longitud de Carga Útil:** Similar al campo de longitud total en IPv4. Especifica la longitud de carga útil, en bytes, que el paquete encapsula.
- e) **Siguiente Encabezado:** El valor de este campo indica el tipo de información que se encuentra a continuación del encabezado IPv6. Esta información puede ser un encabezado TCP o UDP, o un "extensión header", es decir, la información complementaria de capa 3 que se utiliza con propósitos de enrutamiento, seguridad, movilidad entre otros.
- f) **Límite de Salto:** Especifica el máximo número de saltos que un paquete IP puede atravesar. Cada salto o routers disminuye este campo en uno (similar al campo de tiempo de vida [TTL] en IPv4). Como no hay checksum en el encabezado IPv6, el router puede disminuir el campo sin recalcularlo. El recalcular cuesta un valioso tiempo de proceso en los routers con IPv4.
- g) **Dirección de Origen:** Campo de 16 octetos o 128 bits. Identifica el origen del paquete.
- h) **Dirección de Destino:** Campo de 16 octetos o 128 bits. Identifica el destino del paquete.
- i) **Encabezados de Extensión:** Sigue los últimos 8 campos. El número de encabezados de extensión no es fijo, con lo que la longitud total de la cadena del encabezado es de extensión variable.

2.4 Modificación de cabecera IPv4 a IPv6.

- Se eliminaron 6 campos del protocolo IPv4:

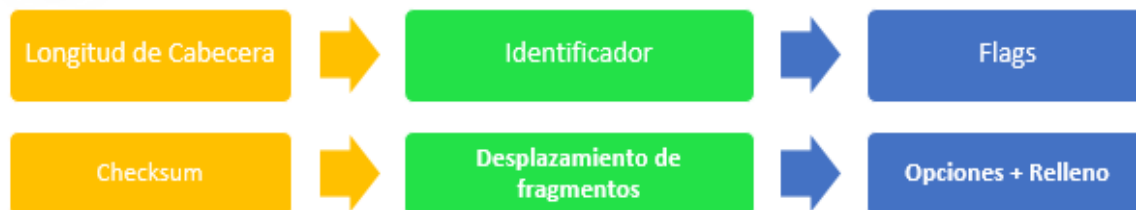


Figura 12: Campos eliminado en IPv6.

- En IPv6 cuatro campos cambiaron el nombre y su ubicación fue modificada.

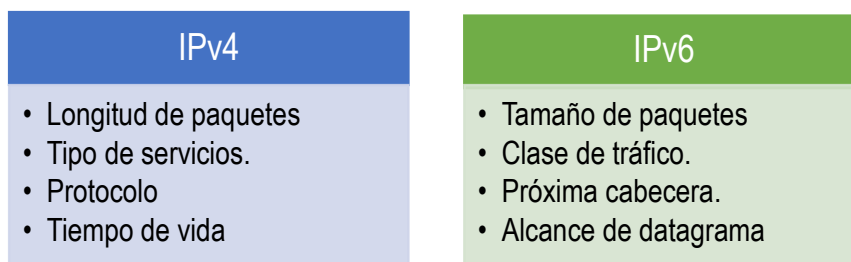


Figura 13: Modificación de campos en IPv6.

- En el protocolo IPv6, el campo identificador de flujo fue agrandado.
- En IPv4 e IPv6 se mantuvieron tres campos

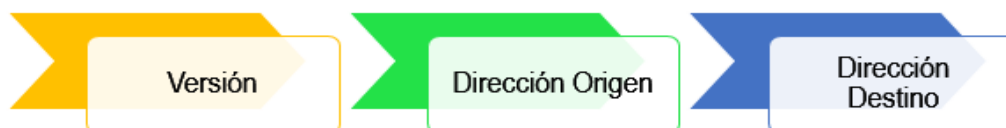


Figura 14: Campos mantenidos en IPv6.

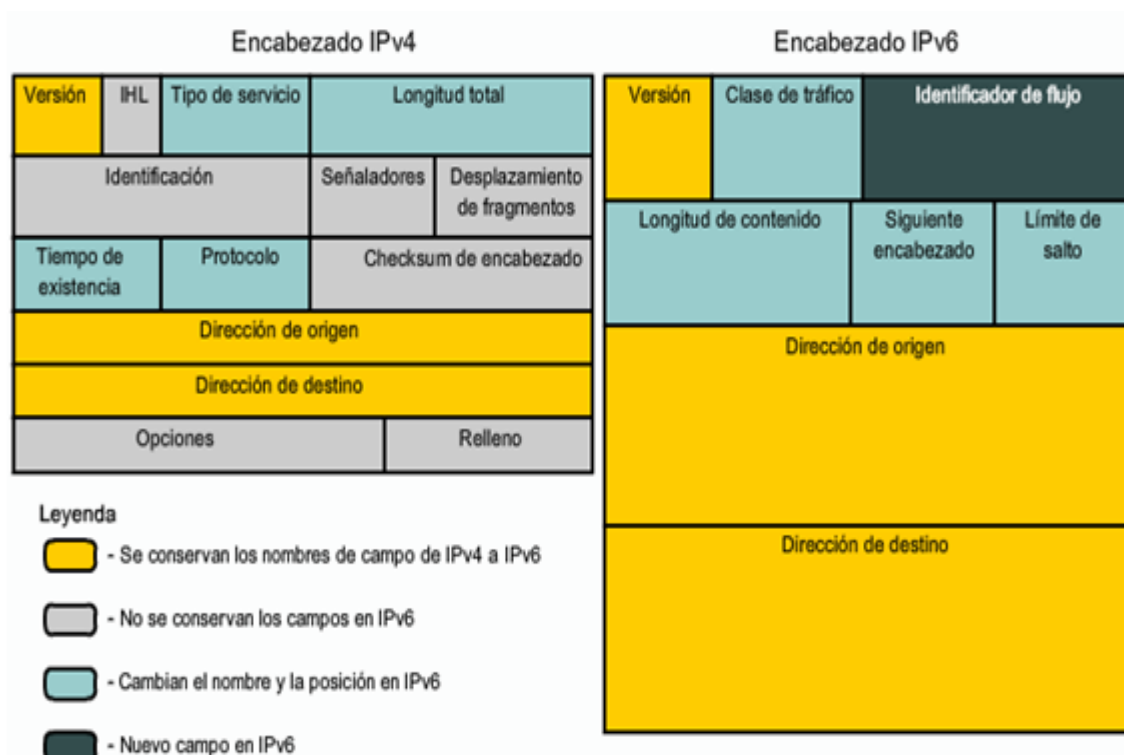


Figura 15: Modificación entre la cabecera IPv4 e IPv6.



2.5 Encabezado de Extensión

En IPv6, la información de capa internet opcional se codifica en cabeceras separadas que se pueden colocar entre la cabecera IPv6 y la cabecera de capa superior dentro de un paquete. Hay un número pequeño de cabeceras de extensión, cada una identificada por un valor. Un paquete IPv6 puede llevar cero, uno, o más cabeceras de extensión, cada una identificada por el campo Cabecera Siguiente de la cabecera precedente.

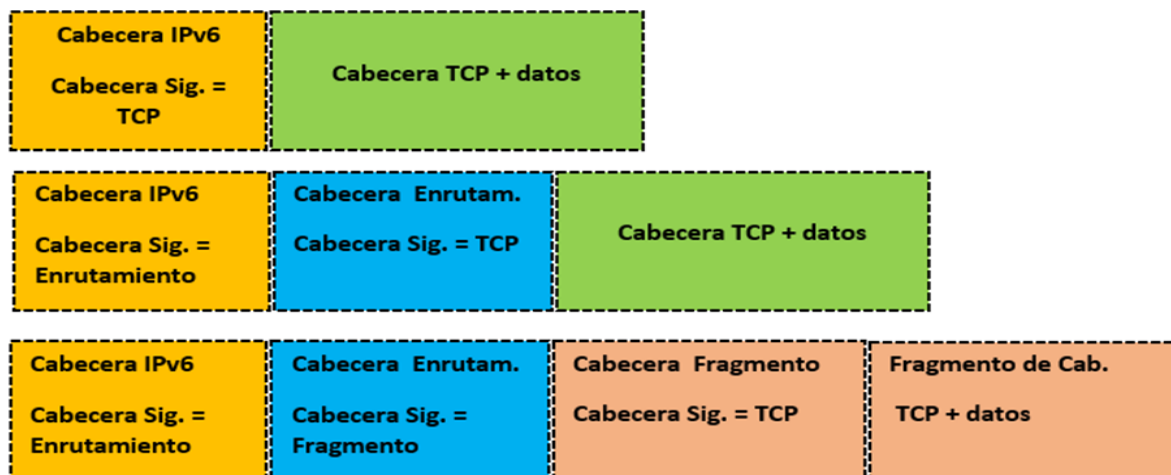


Figura 16: Encabezado de Extensión IPv6.

Con una excepción, las cabeceras de extensión no son examinadas ni procesadas por ningún nodo a lo largo de la ruta de entrega de un paquete, hasta que el paquete alcance el nodo (o cada uno del conjunto de nodos, en el caso de multienvío) identificado en el campo dirección destino de la cabecera IPv6. Allí, el de multiplexaje normal en el campo cabecera siguiente de la cabecera IPv6 invoca el módulo para procesar la primera cabecera de extensión, o la cabecera de capa superior si no hay ninguna cabecera de extensión presente. El contenido y la semántica de cada cabecera de extensión determinan si se procede o no a la cabecera siguiente. Por lo tanto, las cabeceras de extensión se deben procesar estrictamente en el orden que aparecen en el paquete; un receptor no debe, por ejemplo, examinar a través de un paquete buscando un tipo en particular de cabecera de extensión y procesar esa cabecera antes de procesar todas las precedentes.

La excepción mencionada en el párrafo precedente es la cabecera opciones de salto a salto, la cual lleva información que debe ser examinada y procesada por cada nodo a lo largo de la ruta de entrega de un paquete, incluyendo los nodos de origen y de destino. La cabecera Opciones de salto a salto, cuando está presente, debe seguir inmediatamente a la cabecera IPv6. Su presencia es indicada por el valor cero en el campo cabecera siguiente de la cabecera IPv6.

Si como resultado de procesar una cabecera, un nodo necesita proceder a la cabecera siguiente, pero el valor de la cabecera siguiente en la cabecera actual es desconocido por el nodo, debe descartar el paquete



y enviar un mensaje ICMP del problema de parámetro al origen del paquete, con un valor Código ICMP de 1 ("encontrado tipo de cabecera siguiente desconocido") y el campo Puntero ICMP conteniendo el desplazamiento del valor desconocido dentro del paquete original.

La misma acción se debería tomar si un nodo encuentra un valor cabecera siguiente de cero en cualquier cabecera con excepción de una cabecera IPv6. Cada cabecera de extensión es un entero múltiplo de 8 octetos de largo, para conservar la alineación de 8 octetos para las cabeceras subsiguientes. Los campos multiocteto dentro de cada cabecera de extensión se alinean en sus límites naturales, es decir, los campos de ancho de n octetos son colocados en un entero múltiplo de n octetos desde el inicio de la cabecera, para $n = 1, 2, 4$, o 8 .

El encabezado de extensión de IPv6 se clasifica de la siguiente manera:

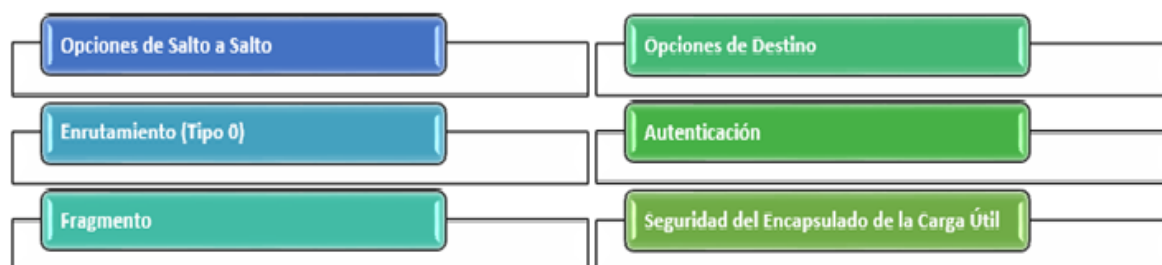


Figura 17: Tipo de encabezado de Extensión IPv6

Cuando más de una cabecera de extensión se usa en un mismo paquete, se debe seguir el siguiente orden:

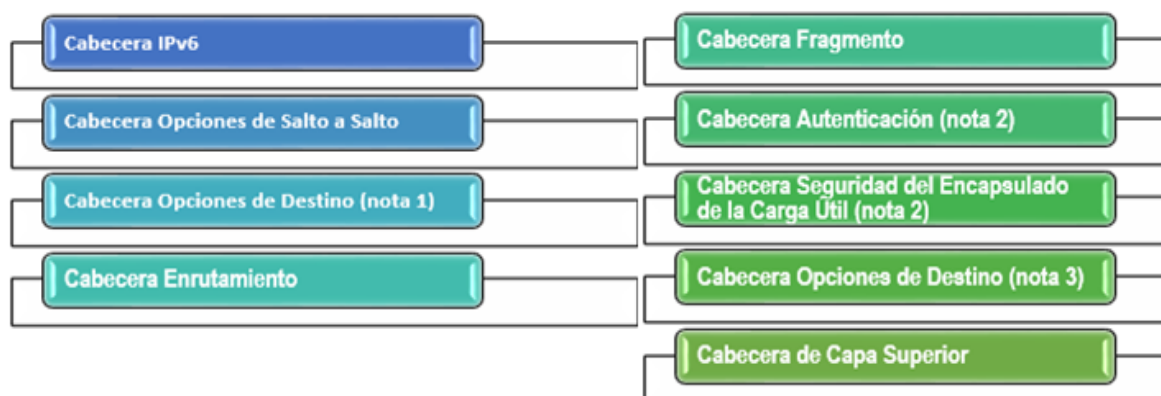


Figura 18: Orden de cabecera de Extensión IPv6.

Nota 1: para las opciones a ser procesadas por el primer destino que aparece en el campo dirección destino IPv6, más los destinos subsiguientes listados en la cabecera enrutamiento.

Nota 2: recomendaciones adicionales con respecto al orden relativo de las cabeceras autenticación y seguridad del Encapsulado de la Carga Útil se dan en la [RFC-2406].

Nota 3: para las opciones a ser procesadas solo por el destino final del paquete.



Cada cabecera de extensión debe ocurrir solamente una vez, a excepción de la cabecera opciones de destino la cual debe de ocurrir a lo sumo dos veces (una vez antes de una cabecera enrutamiento, y la otra antes de una cabecera de capa superior). Si la cabecera de capa superior es otra cabecera IPv6 (en el caso de que el IPv6 sea tunelizado o encapsulado en el IPv6), puede ser seguida por sus propias cabeceras de extensión, las cuales están separadamente sujetas a las mismas recomendaciones de orden. Siempre y cuando se definan otras cabeceras de extensión, sus restricciones de orden concerniente a las cabeceras arriba listadas deben ser especificadas.

Los nodos IPv6 deben aceptar e intentar procesar cabeceras de extensión en cualquier orden y cualquier número de veces que ocurran en un mismo paquete, a excepción de la cabecera Opciones de Salto a Salto la cual está restringida a aparecer sólo inmediatamente después de una cabecera IPv6. No obstante, se aconseja fuertemente que los originadores de paquetes IPv6 se apeguen al orden recomendado arriba, a menos que las especificaciones subsiguientes corrijan esa recomendación.

Opciones

Dos de las cabeceras de extensión actualmente definidas, la **cabecera Opciones de Salto a Salto** y la **cabecera Opciones de Destino** llevan un número variable de "opciones" codificadas de tipo-longitud-valor (TLV), de la siguiente forma:



Figura 19: Formato de cabecera Opción de Salto a Salto y Opción Destino.

- **Tipo de Opción:** Identificador de 8 bits del tipo de opción.
- **Lon-Datos-Opc:** Entero sin signo de 8 bits. Longitud del campo Datos de la Opción.
- **Datos de la Opción:** Campo de longitud variable. Datos específicos del Tipo de Opción

La secuencia de opciones dentro de una cabecera se debe procesar estrictamente en el orden que aparece en la cabecera; un receptor no debe, por ejemplo, examinar a través de una cabecera buscando un tipo en particular de opción y procesar esa opción antes de procesar todas las precedentes.

El identificador **tipo de opción** se codifica internamente tal que sus 2 bits de más alto orden, que especifican la acción que se debe tomar si el nodo IPv6 en proceso no reconoce el tipo de opción:

- 00** - No tomar en cuenta esta opción y continuar procesando la cabecera.
- 01** - Descartar el paquete.



10 - Descartar el paquete y, sin tener en cuenta si o no la dirección destino del paquete fue una dirección multienvío, enviar un mensaje ICMP problema de parámetro, código 2, a la dirección origen del paquete señalando el tipo de opción desconocido.

11 - Descartar el paquete y, solo si la dirección destino del paquete no fue una dirección multienvío, enviar un mensaje ICMP problema de parámetro, Código 2, a la dirección origen del paquete señalando el tipo de opción desconocido.

El tercer bit de más alto orden del tipo de opción especifica si o no los datos de la opción de esa opción pueden modificar el enrutador hacia el destino final del paquete. Cuando una cabecera autenticación está presente en el paquete, para cualquier opción cuyos datos pueden modificar el enrutador, su campo entero datos de la opción se debe tratar como octetos de valor cero cuando se calcula, o verifica el valor de autenticidad del paquete.

0 - Los Datos de la Opción no modifican el enrutador.

1 - Los Datos de la Opción pueden modificar el enrutador.

Los tres bits de alto orden descritos arriba están para ser tratados como parte del tipo de opción, no independientemente del tipo de opción. Es decir, una opción en particular se identifica por un tipo de opción de 8 bits completo, no sólo por los 5 bits de bajo orden de un tipo de opción.

El mismo espacio de enumeración del tipo de opción se usa tanto para la cabecera opciones de salto a salto como para la cabecera opciones de destino. Sin embargo, la especificación de una opción en particular puede restringir su uso a solamente una de esas dos cabeceras.

2.6 Tipos de cabeceras de Extension

Cabecera de Extensión	Tamaño	Descripción
Opciones salto a salto	variable	Contiene datos que deben ser examinados por cada nodo a través de la ruta de envío de un paquete.
Enrutamiento (<i>Routing</i>)	variable	Métodos para especificar la forma de encaminar un datagrama. (Usado con IPv6 Móvil)
Opciones para el destino (<i>Destination Options</i>)	variable	Información que necesita ser examinada solamente por los nodos de destino del paquete
Cabecera de fragmentación	64 bits	Contiene parámetros para la fragmentación de los datagramas.



Cabecera de autenticación	variable	Contiene información para verificar la autenticación de la mayor parte de los datos del paquete
Encapsulado de seguridad de la carga útil	variable	Lleva la información cifrada para comunicación segura.

Tabla 3: Resumen de la Cabecera de Extensión IPv6.

➤ **Cabecera de Extensión Salto a Salto.**

Se utiliza para especificar parámetros de envío en cada salto hacia el destino, **consta de un campo Next Header** (Encabezado siguiente), un campo **Header Extension Length** (Longitud de extensión del encabezado) y un campo **Options** (Opciones) que contiene una o varias opciones.

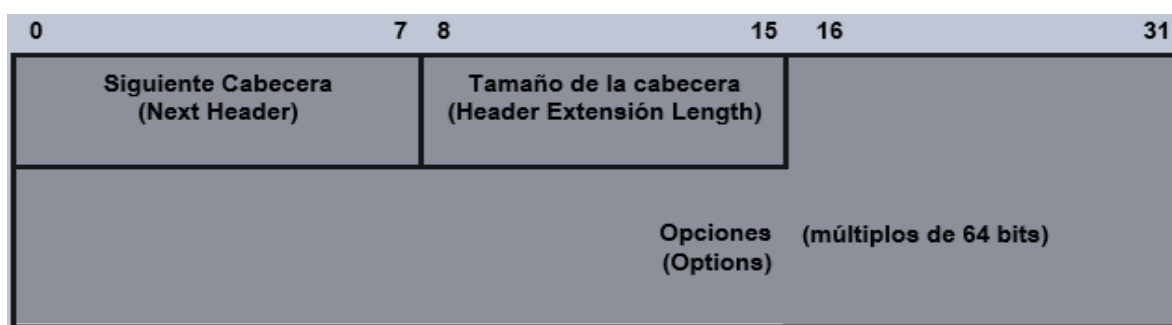


Figura 20: Formato Cabecera de Extensión Salto a Salto.

El valor del campo Header Extension Length es el número de bloques de 8 bytes del encabezado de extensión Salto a Salto, sin incluir los 8 primeros bytes, por lo tanto el valor del campo Header Extension Length es 0. Se utilizan opciones de relleno para garantizar límites de 8 bytes.

Una opción es un encabezado dentro del encabezado de opciones de salto a salto que describe una característica específica de la entrega del paquete o proporciona relleno. Cada opción se codifica en el formato tipo-longitud-valor (TLV), que se utiliza comúnmente en los protocolos TCP/IP. El tipo de opción identifica a la opción y determina el tipo de tratamiento por parte del nodo de procesamiento. La longitud de la opción identifica su longitud. El valor de la opción son los datos asociados a ésta.

➤ **Encabezado Destination Options (Opciones de destino)**

El encabezado Destination Options se utiliza para especificar parámetros de entrega de paquetes para destinos intermedios o para el destino final. Este encabezado se identifica mediante el valor 60 en el campo Next Header (Encabezado siguiente) del encabezado anterior. Los campos del encabezado Opciones de destino se definen del mismo modo que el encabezado Hop-by-Hop Options (Opciones de salto a salto). El encabezado Destination Options se utiliza de dos maneras:



- 1) Si hay un encabezado routing (enrutamiento), especifica opciones de entrega o de proceso en cada destino intermedio.
- 2) También especifica opciones de entrega o de proceso en el destino final.

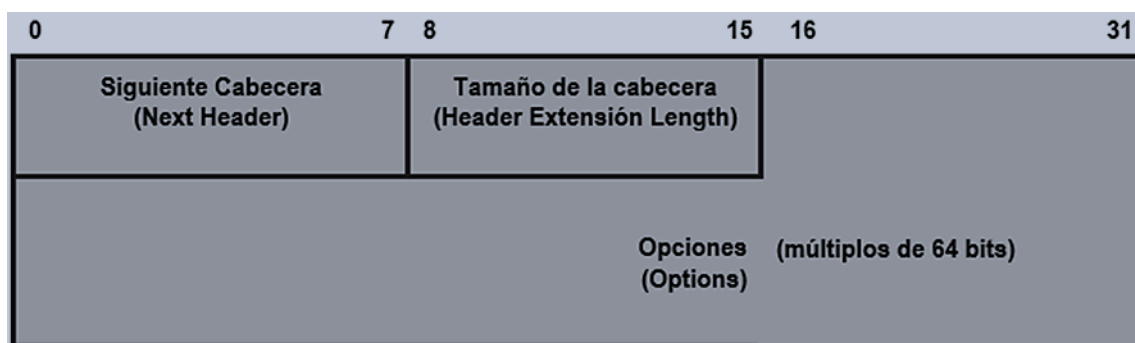


Figura 21: Formato de cabecera de Extensión de Opción Destino.

➤ **Cabecera de Extensión de Enrutamiento.**

De forma similar al enrutamiento de origen que admite IPv4, los nodos de origen de IPv6 pueden utilizar el encabezado de extensión Routing para especificar una ruta de origen, una lista de destinos intermedios para que el paquete viaje por su ruta de acceso al destino final. El encabezado Routing se identifica mediante el valor 43 en el campo Next Header (Encabezado siguiente).

El encabezado Routing consta de un campo Next Header, un campo Header Extensión Length (que se define del mismo modo que en el encabezado de extensión Hop-by-Hop Options), un campo Routing Type (Tipo de enrutamiento), un campo Segments Left (Segmentos restantes) y datos específicos del tipo de enrutamiento.

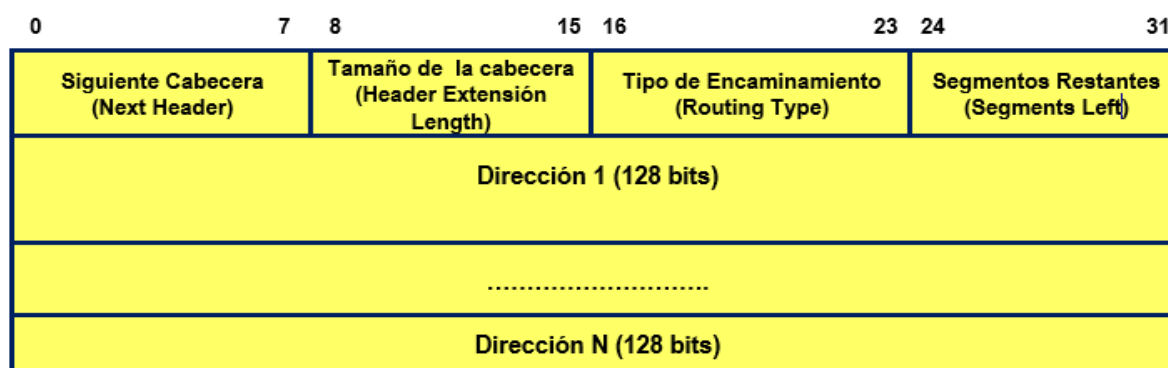


Figura 22: Formato cabecera de Extensión de Enrutamiento



Si, al procesar un paquete recibido, un nodo encuentra una cabecera enrutamiento con un valor tipo de enrutamiento desconocido, el comportamiento requerido del nodo depende del valor del campo segmentos dejados, como sigue:

- Si segmentos dejados es cero, el nodo debe ignorar la cabecera enrutamiento y proceder a procesar la siguiente cabecera en el paquete, cuyo tipo se identifica por el campo cabecera siguiente en la cabecera enrutamiento.
- Si segmentos dejados no es cero, el nodo debe descartar el paquete y enviar un mensaje ICMP problema de parámetro, código 0, a la dirección origen del paquete, apuntando al tipo de enrutamiento desconocido.
- Si, después de procesar una cabecera enrutamiento de un paquete recibido, un nodo intermedio determina que el paquete será remitido hacia un enlace cuya MTU de enlace es menor que el tamaño del paquete, el nodo debe descartar el paquete y enviar un mensaje ICMP "paquete demasiado grande a la dirección origen del paquete".

➤ Cabecera de Extensión de Fragmentación.

Es de gran ayuda en ésta versión del protocolo ya que en la versión anterior no existía un encabezado o un bit de fragmentación y cuando un datagrama llegaba a un router que no lo podía gestionar, éste lo fragmentaba y enviaba el mismo pero fragmentado, si en el camino se perdía alguno de los fragmentos que enviaba el router, se tenía que reenviar todo el datagrama completo, en éstos casos el router generaba más tráfico en la red y no conviene retransmitir todo un datagrama por un fragmento mínimo que no llegó.

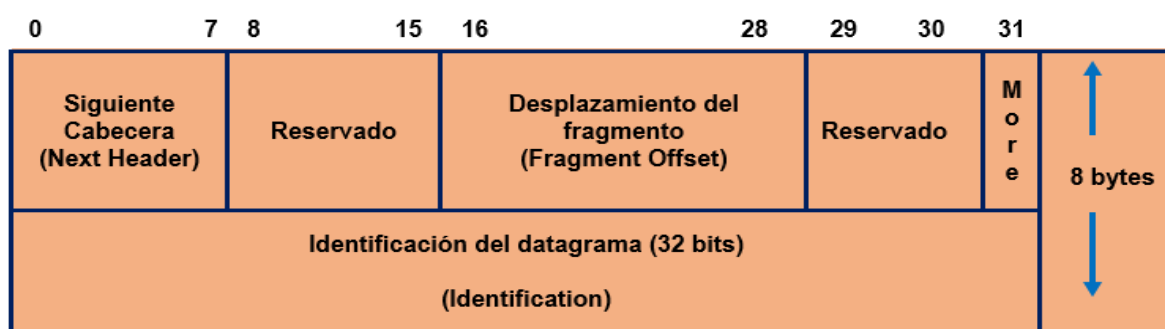


Figura 23: Formato cabecera de Extensión de Fragmentación.

- Next Header:** Especifica el siguiente tipo de encabezado que hay, en el caso de que lo haya.
- Reservado:** 8 bits inicializado a cero para la transmisión; ignorado en la recepción y no se utiliza por el momento.
- Desplazamiento del fragmento:** Indica los 13 bits más significativos del desplazamiento, esto porque la fragmentación se hace en múltiplos de 64.



- d) Los siguientes **2 bits** se han reservado para usos futuros y finalmente el último bit es el más importante ya que indica si hay más fragmentos, si hay más fragmentos su valor es **1** y si no hay más fragmentos que le sigan, su valor es **0**.
- e) **Identificación:** Para enviar un paquete que es demasiado grande para caber en la MTU de la ruta hacia su destino, un nodo origen puede dividir el paquete en fragmentos y enviar cada fragmento como un paquete separado, para ser re ensamblado en el receptor.

En IPv6 sólo el nodo origen puede fragmentar los datos procedentes del nivel superior. Si el tamaño de los datos enviados por el protocolo superior supera la MTU, IPv6 los fragmenta y añade esta cabecera, para permitir el ensamblaje en el destino. Los «routers nunca pueden fragmentar un paquete. Desde el punto de vista de la fragmentación, todo paquete se divide en:

Zona indivisible: Debe ser procesada por los nodos intermedios. Consta de la cabecera IPv6 y las extensiones: «hop-by-hop», opciones de destino intermedio, y enrutamiento.

Zona divisible: Sólo debe ser procesada por el destino, y consta de las extensiones: autenticación, encapsulado de seguridad de la carga, opciones de destino final, y datos del nivel

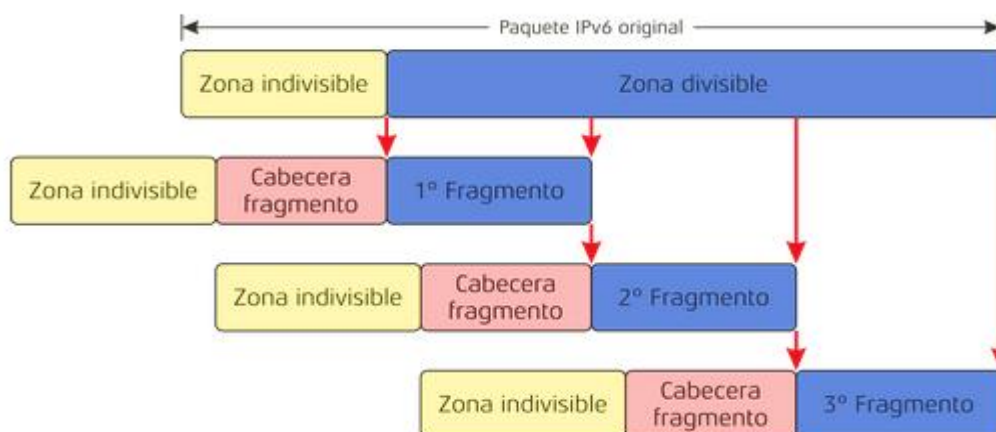


Figura 24: Zona divisible vs indivisible de la fragmentación

➤ Encabezado de Autenticación.

Es una de las mejoras importantes que se mencionaron anteriormente ya que éste encabezado debe estar entre el encabezado IP y los datos del datagrama, no cambia en nada como manejan los datos los protocolos superiores, lo que realiza es proporcionar una seguridad intrínseca en el origen del datagrama, por lo tanto, en cuanto los protocolos de orden superior reciban un datagrama sin la correspondiente autenticación, lo deben desechar.

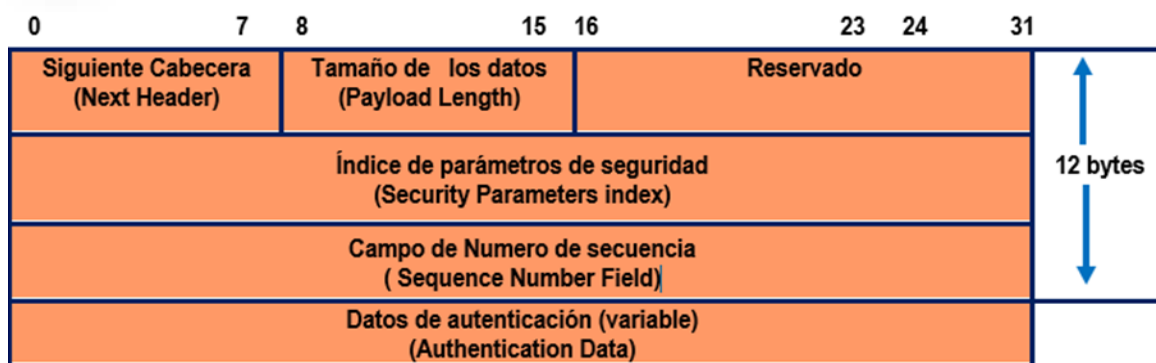


Figura 25: Formato cabecera de Extensión de Fragmentación.

- a) **Siguiente cabecera:** Se pueden encadenar varios encabezados además de éste.
- b) **Longitud (8 bits):** Número de bloques de 4 bytes que componen la cabecera, sin contar los dos primeros.
- c) **Índice de parámetros de seguridad (SPI) (32 bits):** Valor arbitrario que, en combinación con la dirección IP destino y el protocolo de seguridad (AH), identifica de forma única la «asociación de seguridad».
- d) **Los valores en el rango 1 a 255 están reservados por IANA para uso futuro.** El valor 0 está reservado para uso local, específico de la implementación, y no debería aparecer en la red.
- e) **Número de secuencia (32 bits):** Contiene un contador que se incrementa automáticamente. Siempre debe estar presente, incluso aunque el receptor no active el servicio anti-repetición.
- f) **Datos de autenticación:** Campo de longitud variable, que contiene un valor de comprobación de integridad para el paquete. La longitud debe ser un múltiplo de 32 bits.

Para realizar la autenticación se utiliza toda la trama y se quitan los campos que puedan variar (límite de salto en IPv6, se pone a 0 para realizar el cálculo), si se fragmenta la trama la autenticación se realizará extremo a extremo: en origen y destino (después del re ensamblaje). Se aplica una clave criptográfica de al menos 128 bits.



➤ Seguridad del Encapsulado de la Carga Útil.

Proporciona confidencialidad, autenticación e integridad de los datos, y servicios de protección anti-repetición. ESP no proporciona seguridad para la cabecera IPv6, ni para las extensiones situadas delante de la cabecera ESP. El formato de la cabecera es muy diferente a los vistos hasta ahora.

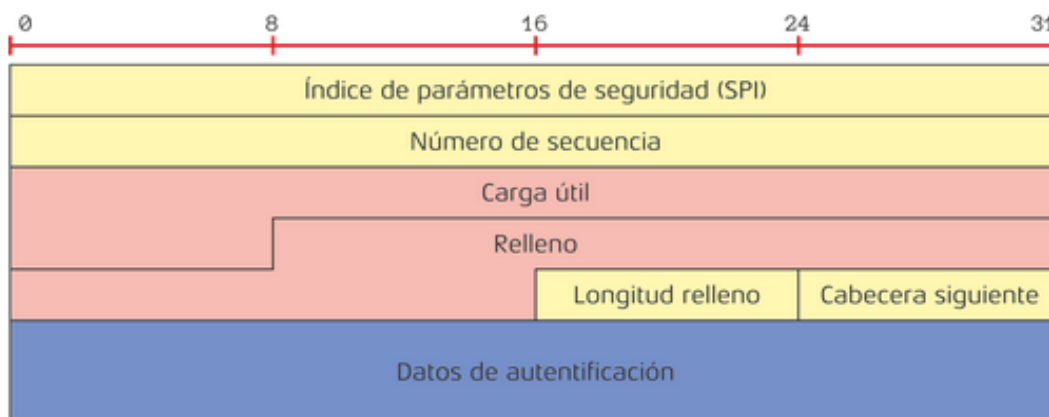


Figura 26: Cabecera ESP

El origen cifra los datos en dos modalidades:

- **Modo transporte.**

Se encripta una parte de la cabecera ESP además del segmento de la capa de transporte.

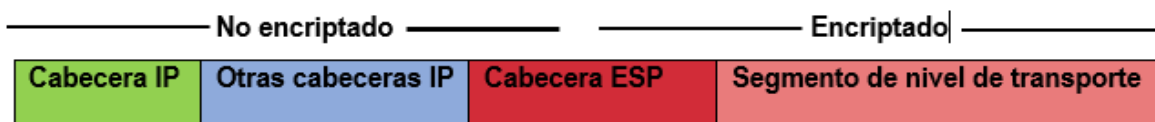


Figura 27: Modo Transporte.

- **Modo túnel.**

Como la cabecera IP ya contiene suficiente información para el encaminamiento la dejamos, pero codificamos todo el paquete IP y parte de la cabecera ESP.

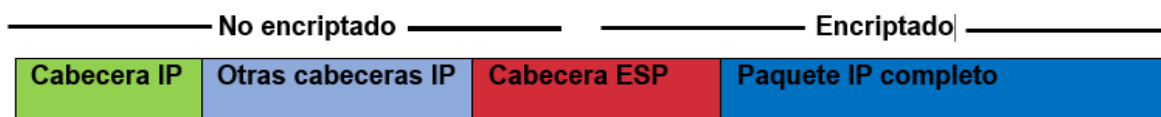


Figura 28: Modo Túnel.



2.7 Ejercicios de análisis de encabezado de extensión IPv6.

1) Comente razonadamente si son correctas o no las siguientes afirmaciones:

- a) En las cabeceras de los datagramas de los protocolos IPv4 e Ipv6 se dispone de un campo de 8 bits cuyo contenido identifica al protocolo receptor de los datos del datagrama.
- b) El tiempo de proceso en los routers para encaminar un datagrama IPv6 será siempre mayor que en el caso de un datagrama IPv4, debido a que la longitud de las direcciones IPv6 son mayores que las direcciones IPv4.
- c) Los routers IPv4 pueden fragmentar datagramas en caso de ser necesarios, pero los routers IPv6 no.
- d) El protocolo IPv6 permite a la máquina de origen fragmentar un datagrama y, por tanto, solo la máquina de destino puede re ensamblarlo. Inicialmente, la máquina de origen necesita recibir de la máquina de destino un mensaje ICMPv6 de paquete demasiado grande indicando la MTU (Unidad Máxima de Transferencia) de su interface de entrada.

2) Conteste razonadamente a las siguientes cuestiones:

- e) Indique, como puede distinguirse en la cabecera IP, tanto en los routers Ipv4 como IPv6, dos comunicaciones: una de transferencias de ficheros y otras de navegación Web entre dos sistemas finales.
- f) Indique, de qué manera puede distinguirse, en el nivel Ip, tanto en los routers IPv4 como IPv6, los paquetes pertenecientes a una comunicación de voz respecto a otra de video (ambas utilizan el protocolo RTP).



SOLUCIÓN

- a) Si. En el caso del protocolo IPv4, pero en IPv6 solo cuando el datagrama no lleve cabeceras de extensión opcionales.
- b) No. El tiempo de proceso es menor en el caso de un datagrama IPv6, debido al menor número de campos de información de control que se ha de examinar.
- c) Si. Los routers IPv6 nunca fragmentan datagramas; en caso de ser necesario debe realizarlo la máquina de origen.
- d) No. Con respecto a que inicialmente la máquina de origen necesita recibir de la máquina de destino un mensaje ICMPv6 de paquete demasiado grande. La máquina de origen solo recibirá este mensaje de los routers intermedios que no soportan la MTU empleada. En IPv6, a diferencia de IPv4, solo la máquina de origen puede fragmentar un datagrama y, por tanto, solo la máquina destinataria (como en IPv4) puede reensamblarlo. Si un router en el camino origen-destino encuentra un datagrama demasiado grande, lo descarta (ya que nunca puede fragmentar) y devuelve al origen un mensaje ICMPv6 de paquete demasiado grande, indicando la MTU del enlace o interfaz de salida. Con esta información, la máquina de origen fragmenta el datagrama en función de la MTU indicada y no como había hecho antes, por omisión, basándose en su MTU de salida.
- e) En IPv4 no se pueden distinguir. En IPv6 puede utilizarse el campo: Etiqueta de Flujo.
- f) En IPv4 no se pueden distinguir. En IPv6 se distinguen mediante la codificación correspondiente de los campos de: Prioridad y Etiqueta de Flujo.



2.8 Notación IPv6.

Representación de direcciones IPv6.

Las direcciones en IPv6 están representadas en la forma x:x:x:x:x:x en donde cada “x” es un fragmento de 16 bits escrito en notación hexadecimal. Cada uno de estos fragmentos debe estar separado por un “:”.

El siguiente es un ejemplo de una dirección IPv6:

3FFE:	8B34:	23C4:	B34A:	023C:	0002:	F436:	1234
16 Bits	16 Bits	16 Bits	16 Bits	16 Bits	16 Bits	16 Bits	16 Bits

Figura 29: Notación IPv6.

Compresión de ceros.

En IPv6 es común que se presenten cadenas grandes de ceros dentro de las direcciones. Para simplificar su escritura se ha convenido en utilizar una sintaxis especial en donde se suprimen los valores consecutivos de ceros ante dos situaciones: campos sucesivos de ceros y campos con ceros al inicio.

Campos con ceros al inicio: Para comprimir direcciones se aplica a cada uno de los campos hexadecimales de 16 bits que tienen uno o más ceros al inicio. Ello involucra que si hay uno o más ceros al inicio de cada campo, estos pueden ser suprimidos para simplificar su longitud y facilitar su lectura y escritura. No obstante, si cada carácter del campo es cero al menos uno debe de ser mantenido.

Ejemplo:

IPv6 Original	3FFE:	8B34 :	<u>023C:</u>	B34A:	<u>003F:</u>	<u>08B3:</u>	23C4 :	<u>0001:</u>
Regla:	3FFE:	8B34 :	23C:	B34A:	3F:	8B3:	23C4 :	1

IPv6 Original	3FFE:	8B34 :	<u>0000 :</u>	<u>0000:</u>	23C4:	<u>0000:</u>	<u>0000:</u>	0001:
Regla:	3FFE:	8B34 :	<u>0:</u>	<u>0:</u>	23C4	<u>0:</u>	<u>0 :</u>	1

Figura 30: Campos con ceros al inicio y ceros totales



Campos sucesivos de ceros: Para simplificar la longitud de una dirección IPv6, cuando se presentan de uno a múltiples campos de ceros, es legal representar estos como ceros o:: (doble dos puntos). Sin embargo, es permitido usarlo una sola vez en la escritura de la dirección.

Ejemplo:

IPv6 Original	3FFE:	8B34 :	0000 :	0000:	0000:	0000:	0000:	0001:
Regla:	3FFE:	8B34	::					1

Figura 31: Campos sucesivos de ceros.

A veces, el mal uso de la compresión de ceros, puede resultar en direcciones erróneas, para ejemplificar lo anterior utilizaremos la siguiente dirección: 3FFE:8B34:0000:0000:3FC0:0000:0000:0001

Cabecera de Extensión	Tamaño	Descripción
No se pueden utilizar dos "::" en una sola dirección. Esto nos llevará a una confusión de cuantos ceros se encuentran entre cada fragmento.	3FFE:8B34::3FC0::1	3FFE:8B34:0000:0000:3FC0::1 ó 3FFE:8B34::3FC0:0000:0000:1
Solamente se pueden omitir los ceros a la izquierda de un fragmento de la dirección. Los ceros a la derecha no pueden ser omitidos	3FFE:8B34:0000:0000:3FC::1	3FFE:8B34:0000:0000:3FC0::1

Tabla 4: Compresión de la notación IPv6

Prefijo IPv6

El aumento de direcciones a 128 bits, garantiza tener innumerables direcciones IP a nuestra disposición, lo que permite tener mejor flexibilidad en la forma en que son asignadas y utilizadas. Al igual que las direcciones IPv4 sin clases, las direcciones IPv6 son divididas en un número de bits pertenecientes al de ID de red seguido por un número de bits pertenecientes al de ID de host.



El prefijo, al igual que en IPv4, indica cuantos bits del lado izquierdo de la dirección identifican la red. Este número debe ser escrito en notación decimal, y debe escribirse al final de la dirección, separado por un “/”, como se muestra en la siguiente dirección:

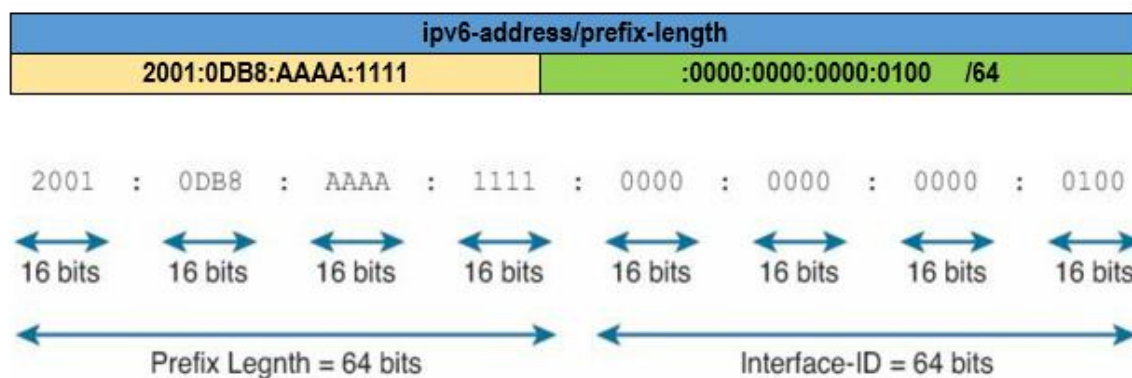


Figura 32: Prefijo IPv6.

El prefijo se puede usar en una dirección en particular como la dirección mostrada previamente, pero también puede ser utilizada para una dirección de red, como la siguiente:

- **2001 : 0DB8 : AAAA : 1111 :: 100 /64**



2.9 Direcciones IPv6

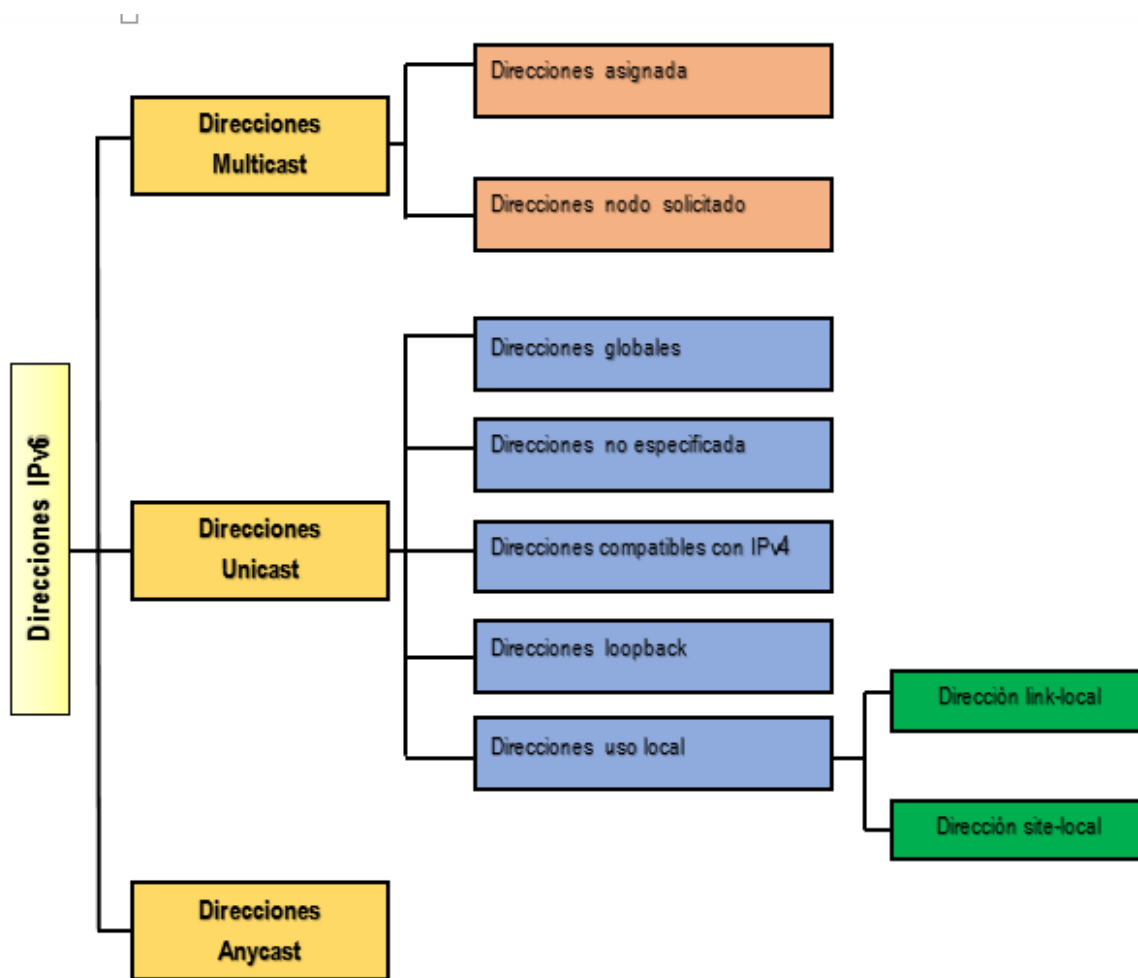


Figura 33: Tipo de direcciones IPv6

Unicast:

- Identifican a un único host en la red.
- Un paquete dirigido a una dirección unicast se entregará únicamente al host identificado con dicha dirección IP.

Multicast:

- Identifican a un grupo de hosts
- Un paquete dirigido a una dirección Multicast se entrega a todos los hosts identificados con esa dirección.
- Implementan también el tráfico Broadcast

**Anycast:**

- Identifican a un grupo de hosts.
- Un paquete dirigido a una dirección anycast se entrega a uno solo de los hosts identificados con esa dirección, normalmente al más cercano, en función de la métrica usada por el protocolo de routing.

2.10 Unicast (Identificación Individual).

Una dirección de tipo unicast está asignada a la mayoría de las veces a una sola interface. Es permitido que varias interfaces tengan la misma dirección unicast, mientras aparezcan como una sola al entorno exterior. Entre las direcciones de este tipo se encuentran las direcciones Globales Agregables, de Uso Local, compatibles con IPv4, así como la dirección de Loopback y la dirección no especificada.

2.11 Tipos de direcciones Unicast

➤ **Dirección no Especificada (Dirección 0:0:0:0:0:0:0:0, o simplemente::).**

No es asignada a ningún nodo. Se emplea para indicar la ausencia de dirección; por ejemplo, cuando se halla en el campo de dirección fuente, indica que se trata de un host que está iniciándose, antes de que haya aprendido su propia dirección.

➤ **Dirección Loopback: Loopback o Dirección de auto-retorno (::1).**

No ha de ser asignada a una interfaz física ya que se trata de una interfaz "virtual" (paquetes que no salen de la máquina que los emite) y que nos permite hacer un bucle para verificar la correcta inicialización del protocolo dentro de una determinada máquina.

➤ **Direcciones Unicast Globales.**

En la actualidad ya se emplea este tipo de direcciones, basadas en la agregación por parte de los proveedores de Internet, y los mecanismos adoptados para IPv6, permiten su continuidad. Pero además, se incorporó un mecanismo de agregación basado en "intercambios". La combinación de ambos es la que permite un encaminamiento mucho más eficiente, dando dos opciones de conectividad a unas u otras entidades de agregación.

Se trata de una organización basada en tres niveles:

1. Topología Pública: conjunto de proveedores e "intercambiadores" que proporcionan servicios públicos de tránsito Internet.
2. Topología de Sitio: redes de organizaciones que no proporcionan servicios públicos de tránsito a nodos fuera de su propio "sitio".
3. Identificador de Interfaz: identifican interfaces de enlaces.



El formato de las direcciones Unicast Globales Agregables es el siguiente:

3	13	8	24	16	64 bits
FP	TLA ID	RES.	NLA ID	SLA ID	INTERFAZ ID
Topología Pública			Topología de Sitio		Identificador de Interfaz

FP	Prefijo de formato
TLA ID	Identificador de agregación de Nivel Superior
RES.	Reservado para uso futuro
NLA ID	Identificador de agregación de siguiente Nivel
SLA ID	Identificador de agregación de Nivel de Sitio
Interfaz ID	Identificador de interfaz

Figura 34: Direcciones Unicast Globales Agregables

TLA ID (identificador de Agregación de Nivel Superior): Se trata del nivel superior en la estructura jerárquica de enrutador. Los routers situados en este nivel tienen, en la tabla de encaminado, una entrada para cada TLA ID activo, y probablemente entradas adicionales relativas al propio TLA ID donde están físicamente situados. Podrían tener otras entradas, para su optimización, dependiendo de su topología, pero siempre pensando en que se minimice la tabla.

RES: El campo Reservado (RES): permitirá, en el futuro, ampliaciones "organizadas" del protocolo, por ejemplo ampliar el número de bits de los campos TLA y NLA. Por el momento contiene ceros.

NLA ID. (Identificador de Agregación de Siguiente Nivel): Es empleado por organizaciones a las que se ha asignado un TLA, para crear una estructura jerárquica de direccionamiento, acorde con su propia red, y para identificar los "sitios" u organizaciones que de ella dependen. Pueden reservar los bits superiores para la diferenciación de la estructura de su red, en función a sus propias necesidades.

SLA. (Identificador de Agregación de Nivel de Sitio): El SLA es usado por organizaciones "finales" para crear su propia estructura jerárquica de direcciones e identificar sus subredes. Es equivalente al concepto de subred en IPv4, con la muy apreciable diferencia de que cada corporación tiene un mayor número de subredes (16 bits proporcionan capacidad para 65.535).



➤ **Dirección Unicast de uso Local.**

Los nodos IPv6 pueden no tener ningún conocimiento o mínimo de la estructura interna de las direcciones IPv6, dependiendo de su misión en la red (por ejemplo, host frente a routers). Pero como mínimo, un nodo debe considerar que las direcciones Unicast (incluyendo la propia), no tienen estructura:

128 bits
Dirección del nodo

Un host algo más sofisticado, conocería el prefijo de la subred del enlace al que esta conectado:

n bits	128-n bits
Prefijo de subred	Identificador de interfaz

Figura 35: Estructura Uso Local.

El "identificador de interfaz" se emplea, por tanto, para identificar interfaces en un enlace, y deben de ser únicos en dicho enlace. En muchos casos también serán únicos en un ámbito más amplio. El mismo identificador de interfaz puede ser empleado en múltiples interfaces del mismo nodo, sin afectar a su exclusividad global en el ámbito IPv6.

Se han definido dos tipos de direcciones unicast de uso local: **Local de Enlace (Link-Local)** y **Local de Sitio (Site-Local)**.

Direcciones locales de enlace: Han sido diseñadas para direccionar un único enlace para propósitos de auto-configuración (mediante identificadores de interfaz), descubrimiento del vecindario, o situaciones en las que no hay routers. Por tanto, los encaminadores no pueden retransmitir ningún paquete con direcciones fuente o destino que sean locales de enlace (su ámbito está limitado a la red local). Tienen el siguiente formato:

FE80::<ID de interfaz>/10.

10 bits	54 bits	64 bits
1111111010	0	Identificador de interfaz

Figura 36: Dirección Link-local

Las direcciones locales de sitio: Permiten direccionar dentro de un "sitio" local u organización, sin la necesidad de un prefijo global. Se configuran mediante un identificador de subred, de 16 bits.



Los encaminadores no deben de retransmitir fuera del sitio ningún paquete cuya dirección fuente o destino sea "local de sitio" (su ámbito está limitado a la red local o de la organización).

FECO::<ID de subred>:<ID de interfaz>/10.

10 bits	38 bits	16 bits	64 bits
1111111011	0	ID de subred	Identificador de interfaz

Figura 37: Dirección Site-local

➤ Direcciones IPv6 con Direcciones IPv4.

Algunas direcciones IPv6 contienen direcciones IPv4 dentro de ellas, y existen dos tipos. Las "Direcciones IPv6 compatibles con IPv4" son para nodos que soportan ambos protocolos (IPv6 e IPv4), y se conforman de la siguiente manera:

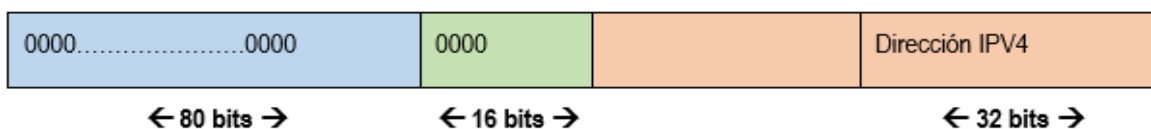


Figura 38: Direcciones IPv6 con direcciones IPv4.

Las "Direcciones IPv6 mapeadas a IPv4" son para nodos que solamente soportan IPv4, y se conforman de la siguiente manera:

2.12 Anycast (Identificación Selectiva).



Figura 39: Dirección IPv6 mapeada a IPv4.

Las direcciones anycast tienen el mismo rango de direcciones que las unicast. Cuando una dirección unicast es asignada a más de una interfaz, convirtiéndose en una dirección anycast, los nodos a los que dicha dirección ha sido asignada, deben ser explícitamente configurados para que reconozcan que se trata de una dirección anycast.



Existe una dirección anycast, requerida para cada subred, que se denomina "dirección anycast del routers de la subred" (Subnet-router anycast Address). Su sintaxis es equivalente al prefijo que especifica el enlace correspondiente de la dirección unicast, siendo el indicador de interfaz igual a cero

n bits	128-n bits
Prefijo de subred	000000000000000000000000

Figura 40: Anycast con indicador de interfaz igual a cero

El prefijo de subred en una dirección anycast es el prefijo que identifica un enlace específico. Esta dirección anycast es sintácticamente igual a una dirección unicast para una interface en el enlace con el ID de interface puesto en cero. Los paquetes enviados a la dirección anycast Subnet-Router serán entregados a un routers en la subred:

- Todos los routers deben soportar las direcciones anycast Subnet-Router para las subredes a las cuales tengan interfaces.
- La dirección anycast subnet-router está diseñada para ser utilizada en aplicaciones donde un nodo necesita comunicarse con alguno del conjunto de routers.

Dentro de cada subred, los 128 valores superiores de identificadores de interfaz están reservados para su asignación como direcciones anycast de la subred.

La construcción de una dirección reservada de anycast de subred depende del tipo de direcciones IPv6 usadas dentro de la subred. Las direcciones cuyos tres primeros bits (prefijo de formato) tienen valores entre 001 y 111 (excepto las de Multicast, 1111 1111), indican con el bit "universal/local" igual a cero, que el identificador de interfaz tiene 64 bits, y por tanto no es globalmente único (es local). En este caso, las direcciones reservadas anycast de subred se construyen del siguiente modo:



64 bits	57 bits	7 bits
Prefijo de subred	1111110111 ... 111	ID anycast
Identificador de Interfaz		

En el resto de los casos, el identificador de interfaz puede tener una longitud diferente de 64 bits, por lo que la construcción se realiza según el siguiente esquema:

n bits	121-n bits	7 bits
Prefijo de subred	1111111...1111111	ID anycast
Identificador de Interfaz		

Figura 41: Ejemplo de direcciones Anycast.

2.13 Multicast

Una dirección Multicast en IPv6, puede definirse como un identificador para un grupo de nodos. Un nodo puede pertenecer a uno o varios grupos Multicast. Las direcciones Multicast tienen el siguiente formato:

8	4	4	112 bits
11111111	000T	Ámbito	Identificador de Grupo

Figura 42: Multicast

El formato del prefijo de una dirección Multicast es 1111 1111, lo cual facilita su identificación, ya que todas las direcciones de este tipo comienzan con 0xFF. Las banderas son 4 bits, aunque en realidad 3 son reservados y deben estar siempre en 0. Solamente se usa el bit de más a la derecha que es conocido como el bit "T". Cuando este bit se encuentra en 0, indica una dirección Multicast que está permanentemente asignada ("Bien conocidas"), y es asignada por una autoridad de numeración en Internet. Cuando este bit se encuentra en 1, indica una dirección Multicast no permanentemente asignada. El campo de alcance es de 4 bits, y sirve para restringir el tráfico de Multicast.



A continuación se muestran los posibles valores:

0	Reservado	9	No Asignado
1	Ámbito Local de Nodo	A	No Asignado
2	Ámbito Local de Enlace	B	No Asignado
3	No Asignado	C	No Asignado
4	No Asignado	D	No Asignado
5	Ámbito Local de Sitio	E	Ámbito Global
6	No Asignado	F	Reservado
7	No Asignado		
8	Ámbito Local de Organización		

Tabla 5: Formato del prefijo de una dirección Multicast.

El "identificador de Grupo", identifica, como cabe esperar, el grupo de Multicast concreto al que nos referimos, bien sea permanente o temporal, dentro de un determinado ámbito.

Por ejemplo, si asignamos una dirección Multicast permanente, con el identificador de grupo 101 (hexadecimal), al grupo de los servidores de tiempo (NTS), entonces:

- **FF01::101** significa todos los NTS en el mismo nodo que el paquete origen.
- **FF02::101** significa todos los NTS en el mismo enlace que el paquete origen.
- **FF05::101** significa todos los NTS en el mismo sitio que el paquete origen.
- **FF0E::101** significa todos los NTS en Internet.

Las direcciones Multicast no permanentes, sólo tienen sentido en su propio ámbito. Por ejemplo, un grupo identificado por la dirección temporal Multicast local de sitio FF15::101, no tiene ninguna relación con un grupo usando la misma dirección en otro sitio, ni con otro grupo temporal que use el mismo identificador de grupo (en otro ámbito), ni con un grupo permanente con el mismo identificador de grupo.

Las direcciones Multicast no deben ser usadas como dirección fuente en un paquete IPv6, ni aparecer en ninguna cabecera de encaminado.



Las principales direcciones Multicast reservadas son las incluidas en el rango FF0x:0:0:0:0:0:0. Algunos ejemplos útiles de direcciones Multicast, según su ámbito, serían:

- FF01:0:0:0:0:0:1** - todos los nodos (ámbito local)
- FF02:0:0:0:0:0:1** - todos los nodos (ámbito de enlace)
- FF01:0:0:0:0:0:2** - todos los routers (ámbito local)
- FF02:0:0:0:0:0:2** - todos los routers (ámbito de enlace)
- FF05:0:0:0:0:0:2** - todos los routers (ámbito de sitio)

La dirección FF02:0:0:0:1:FFxx:xxxx, denominada "Solicited-Node Address", o dirección de nodo solicitada, permite calcular la dirección Multicast a partir de la unicast o anycast de un determinado nodo. Para ello, se sustituyen los 24 bits de menor peso ("x") por los mismos bits de la dirección original. Así, la dirección 4037::01;800:200E:8C6C se convertiría en FF02::1:FF0E:8C6C.

Cada nodo debe de calcular y unirse a todas las direcciones Multicast que le corresponden para cada dirección unicast y anycast que tiene asignada.

2.14 Plan de direccionamiento.

Direcciones obligatorias en un Host IPv6:

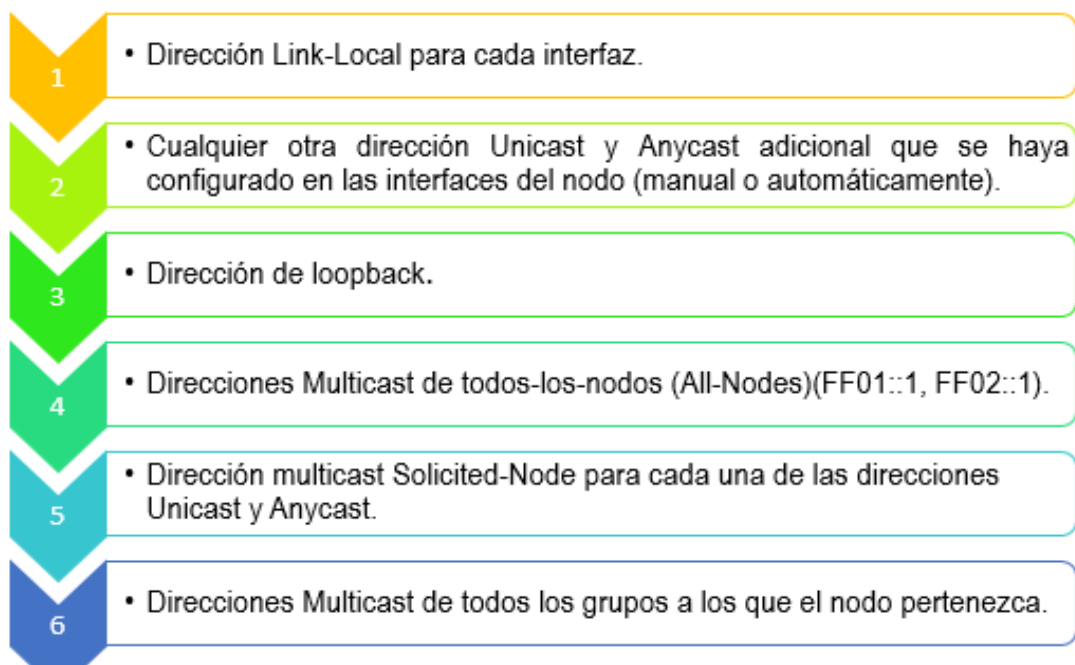


Figura 43: Direcciones obligatorias en un host IPv6.

Direcciones obligatorias en un routers IPv6:

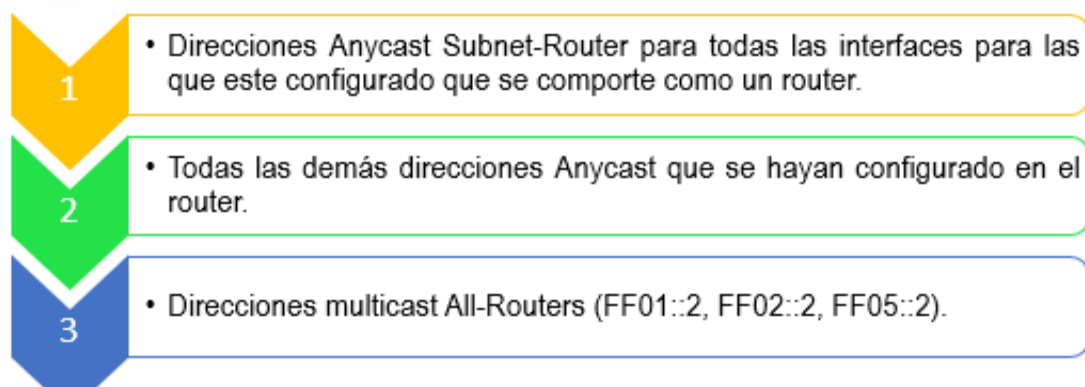


Figura 44: Direcciones obligatorias en un router IPv6

El plan de direccionamiento o numeración tiene como objetivo la asignación del direccionamiento IPv6. Dicha asignación es para las diferentes redes y subredes existentes en una red determinada, para ello se pueden considerar los siguientes criterios (RFC3177 y tendencias reales):

Todas las redes internas que vayan a desplegar IPv6 tendrán un prefijo /64, necesario para la construcción automática de direcciones IPv6 de tipo unicast y/o Anycast.

Los usuarios finales, clientes residenciales (acceso xDSL, FTTx, etc.), como corporativos (empresas, ISPs, Universidad, etc.) podrán recibir prefijos de longitud /48, y posibilita crear hasta 216 (65.536) subredes IPv6 de prefijo /64.

Para la elaboración del plan de direccionamiento se deben tener en cuenta las diversas subredes existentes susceptibles de desplegar IPv6 en algún momento, éstas pueden incluir:

- Subredes susceptibles de ser nativas IPv6 desde el primer momento del despliegue de IPv6.
- Subredes susceptibles de ser nativas IPv6 a medio o largo plazo, no necesariamente desde el comienzo del despliegue de IPv6.
- Servicios de transición a IPv6.

2.16 Subnetting en IPv6

En el subnetting del protocolo IPv6, necesitamos recordar los siguientes aspectos:

- Cada carácter de una dirección IPv6 representa un nibble (4 bits).
- Se compone de ocho grupos de caracteres hexadecimales (los números 0–9 y las letras A–H) separados por dos puntos;
- Una vez que se hace en binario nada cambia. Es fácil perderse en tantos bits, pero la matemática es toda igual. Cada bit de subred es uno menos de host y viceversa.

**Ejemplo:**

Una Empresa necesita crear 65, 536 Subredes, con el siguiente bloque de dirección: 2001:0DB8:ACAD::/48

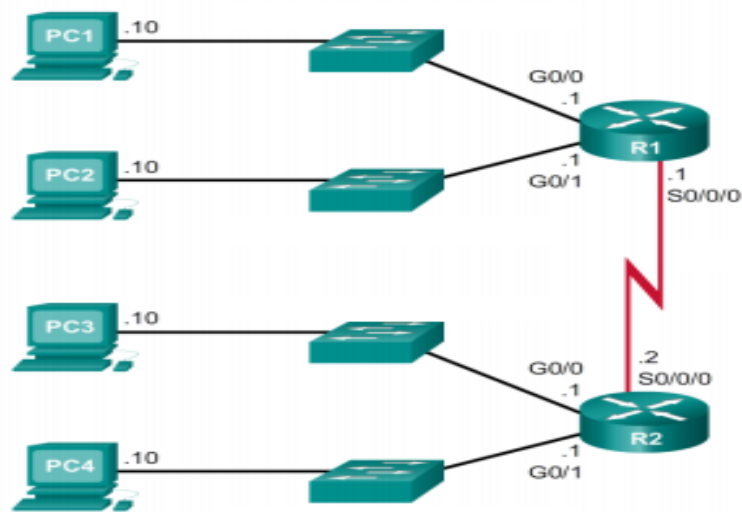


Figura 45: Topología de empresa.

Mostrar la longitud exacta de la dirección de red:

2001:	0db8:	Acad	0000	0000	0000	0000	0000	/48
-------	-------	------	------	------	------	------	------	-----

Figura 46: Descomposición de dirección IPv6 de la Empresa.

La empresa a necesitar 65,536 subredes, es requerido 16 bit, 2^{16} , por lo cual se le suma a las 48 bit que tenía anteriormente, para pasar de un (/48) a (/64).

Descomponemos la dirección de red, en binario requeridos, para realizar la combinación determinada:

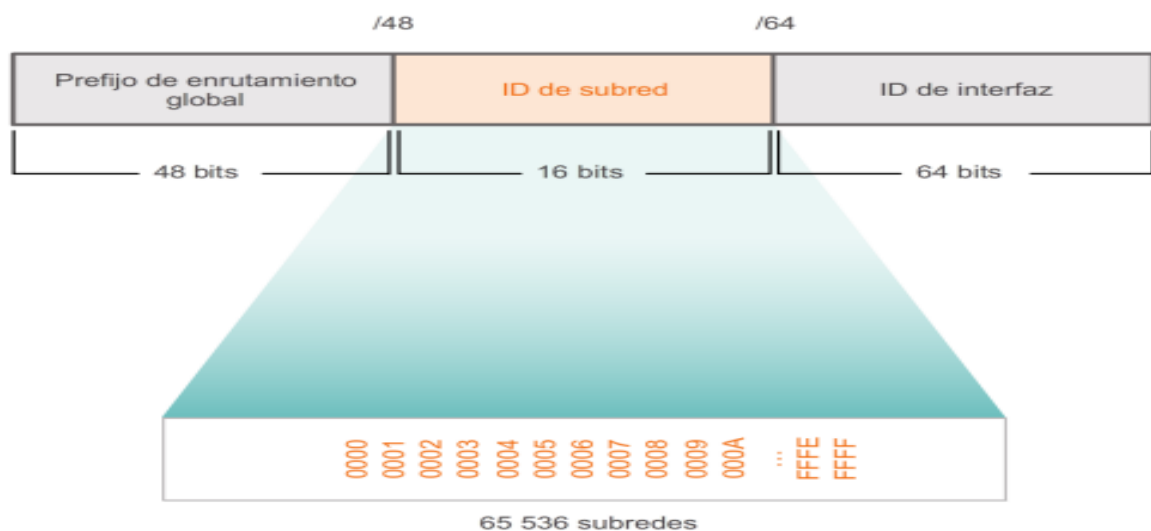


Figura 47: Combinación de binario requeridos para subnetear dirección IPv6



Una vez realizado la combinación de binarios podremos obtener cada una de las siguientes subredes:



Figura 48: Separación de los bit de red y bit de host

- Subred 1: 2001:0DB8:ACAD:0001::/64
- Subred 2: 2001:0DB8:ACAD:0002::/64
- Subred 3: 2001:0DB8:ACAD:0003::/64
- Subred 4: 2001:0DB8:ACAD:0004::/64.
- Sucesivamente.

Esquema Final:

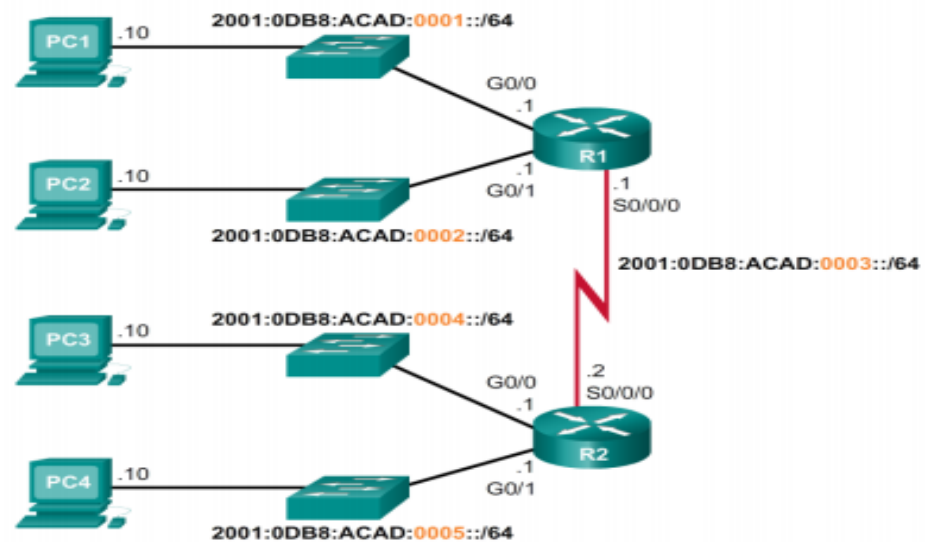


Figura 49: Esquema final de la empresa.



2.17 Ejercicio de IPv6.

1. **¿Cuál de los siguientes es la abreviatura válida más corta para FE80: 0000: 0000: 0100: 0000: 0000: 0000: 0123?**
 - a) FE80::100::123
 - b) FE8::1::123
 - c) FE80::100:0:0:0:123:4567
 - d) FE80:0:0:100::123

2. **¿Cuál de las siguientes es la abreviatura válida más corta para la siguiente dirección 2000: 0300: 0040: 0005: 6000: 0700: 0080: 0009?**
 - a) 2:3:4:5:6:7:8:9
 - b) 2000:300:40:5:6000:700:80:9
 - c) 2000:300:4:5:6000:700:8:9

3. **¿Cuál de los siguientes es el prefijo de la dirección 2000: 0000: 0000: 0005: 6000: 0700: 0080: 0009, en el supuesto de una máscara de / 64?**
 - a) 2000:3:4:5:6:7:8:9
 - b) 2000::5::/64
 - c) 2000::5:0:0:0:0/64
 - d) 2000:0:0:5::/64
 - e) 2000:0:0:5:0:0:0:0/64

4. **¿Cuál de las siguientes direcciones de multidifusión se define como la dirección para el envío de paquetes sólo a los routers IPv6 del vínculo local?**
 - a) FF02::1
 - b) FF02::2
 - c) FF02::5
 - d) FF02::A

5. **Usted tiene un prefijo / 32 a partir de 2001 0db8. ¿Cómo se puede buscar en la base de datos RIPE?**
 - a) 2001:0db8
 - b) 2001:0db8/32
 - c) 2001:0db8::/32
 - d) 2001:db8::/32



6. ¿Cómo se puede comprimir correctamente la dirección IPv6 a seguir 2001: 0db8: 0000: 0000: 0000: 0000: 0000: 0C50?
- a) 2001:0db8:0:0:0:0:0:0c50
 - b) 2001:0db8::0c50
 - c) 2001:db8::c50
 - d) 2001:db8::c5
7. ¿Cómo se puede comprimir correctamente la dirección IPv6 a seguir 2001: 0db8: 0000: 0000: b450: 0000: 0000: 00b4?
- a) 2001:db8::b450::b4
 - b) 2001:db8::b450:0:0:b4
 - c) 2001:db8::b45:0000:0000:b4
 - d) 2001:db8:0:0:b450::b4
8. ¿Cómo se puede comprimir correctamente la dirección IPv6 a seguir: 2001 0db8: 00F0: 0000: 0000: 03d0: 0000: 00ff?
- a) 2001:0db8:00f0::3d0:0:00ff
 - b) 2001:db8:f0:0:0:3d0:0:ff
 - c) 2001:db8:f0::3d0:0:ff
 - d) 2001:0db8:0f0:0:0:3d0:0:0ff
9. ¿Cómo se puede acceder a su servidor web en IPv6 2001: db8 :: 8080 en el puerto 8080 utilizando un navegador web ?
- a) http://2001:db8::8080:8080
 - b) http://2001:0db8:0000:0000:0000:0000:8080:8080
 - c) http://[2001:db8::8080]:8080
 - d) You cannot use the IPv6 address, you have to rely on DNS.
10. ¿Cómo se puede comprimir correctamente la dirección IPv6 a seguir: 2001 0db8: 0f3c: 00d7: 7dab: 03d0: 0000: 00FF?
- a) 2001:db8:f3c:d7:7dab:3d:0:ff
 - b) 2001:db8:f3c:d7:7dab:3d0:0:ff
 - c) 2001:db8:f3c:d7:7dab:3d0::ff
 - d) 2001:0db8:0f3c:00d7:7dab:03d::00ff



3. ICMPv6 (Control Internet Message Protocol version 6)

Dado que IPv6 es una versión más reciente, el protocolo ICMP fue mejorado de manera significativa y se enfatiza en las siguientes funciones:

- Asume el papel de algunos protocolos auxiliares en IPv4: IGMP (Protocolo de gestión de grupos Multicast) y ARP (protocolo de resolución de direcciones).
- Realiza Diagnóstico (Ping).
- Descubrimiento de routers vecinos.
- Autoconfiguración de interfaz.
- Resolución de direcciones (IP a capa de enlace).
- Detección de próximo salto (ruta por defecto).
- Detección de caídas de vecinos.
- Detección de direcciones duplicadas.
- Redireccionamiento.

ICMPv6 es un protocolo orientado a mensajes:

- Mensajes de error.
- Mensajes de información.
- Mensajes para el descubrimiento de vecinos.
- Mensajes de pertenencia y gestión de grupos.

3.1 Formato de mensajes ICMPv6

Todos los mensajes ICMPv6 tienen un formato común:

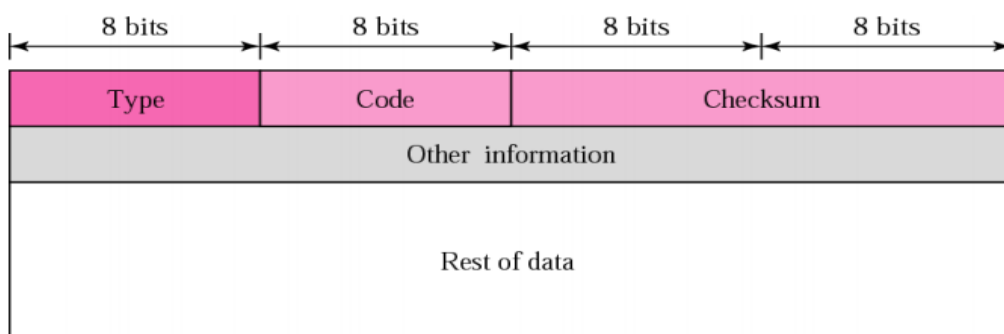


Tabla 6: Formato de mensajes ICMPv6



- a) **El campo “tipo”** indica el tipo de mensaje, y su valor determina el formato del resto de la cabecera.
- Si el bit de mayor peso tiene el valor 0 (valores entre 0 y 127) entonces es un mensaje de error.
 - Si el bit de mayor peso es 1 (valores entre 128 y 255) entonces es un mensaje informativo.
- b) **El campo Código (8bits)** depende del tipo de mensaje, y son usados para crear un nivel adicional de clasificación de mensajes, de tal forma que los mensajes informativos en función del campo código se pueden subdividir en varios tipos.
- c) **El campo Checksum** nos permite detectar errores en el mensaje ICMPv6.

3.2 Mensajes de error ICMPv6

Los mensajes de error de ICMPv6 son similares a los mensajes de error de ICMPv4. Se dividen en 4 categorías: destino inaccesible, paquete demasiado grande, tiempo excedido y problemas de parámetros.

- **Mensaje de destino inalcanzable:** Se genera un mensaje de destino inalcanzable (tipo 1), en respuesta a un paquete que no puede entregarse a su dirección de destino por razones distintas a la congestión. Las razones de la falta de entrega de un paquete son descritas por el valor de .campo del código.
- **Mensaje de paquete demasiado grande:** se envía en respuesta a un paquete que no pueda remitir, porque el paquete es más grande que la unidad máxima de transmisión (MTU) (MTU) del link saliente.
- **Tiempo Excedido:** Si un router recibe un paquete con un límite de salto de cero, o si un router reduce el límite de salto de un paquete a cero, debe descartar el paquete y enviar un mensaje de tiempo excedido ICMPv6 con código 0 al origen del paquete. Esto indica un Routing Loop o un valor límite inicial del salto que sea demasiado pequeño.
- **Problemas de Parámetros:** Se genera en respuesta a un paquete del IPv6 con el problema en su encabezado del IPv6, o los encabezados de extensión, nodo no puede procesar el paquete y debe desecharlo.

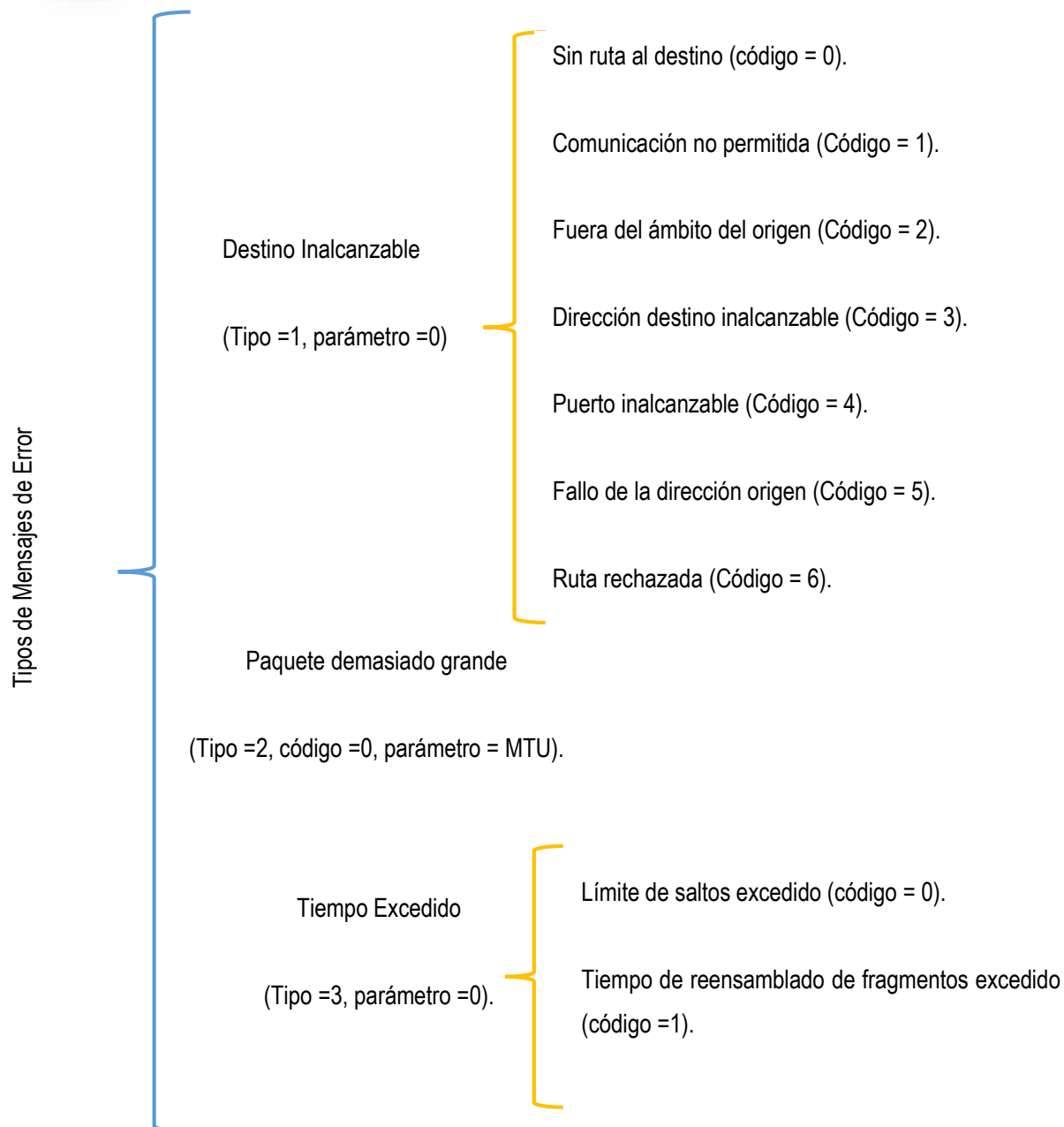


Figura 50: Mensajes de error de ICMPv6



3.3 Ejemplo Destino Inalcanzable

- Un router o host no pueden encaminar o entregar un datagrama.
- El datagrama se descarta y se envía un mensaje ICMP de error.

Type: 1	Code: 0 to 4	Checksum
Unused (All 0s)		
Part of the received IP datagram including IP header plus the first 8 bytes of datagram data		

Figura 51: Ejemplo Destino Inalcanzable

3.4 Mensajes de Información

- Incluye los mensajes de echo-request y echo-reply.
- Permite identificar errores en la capa de red.

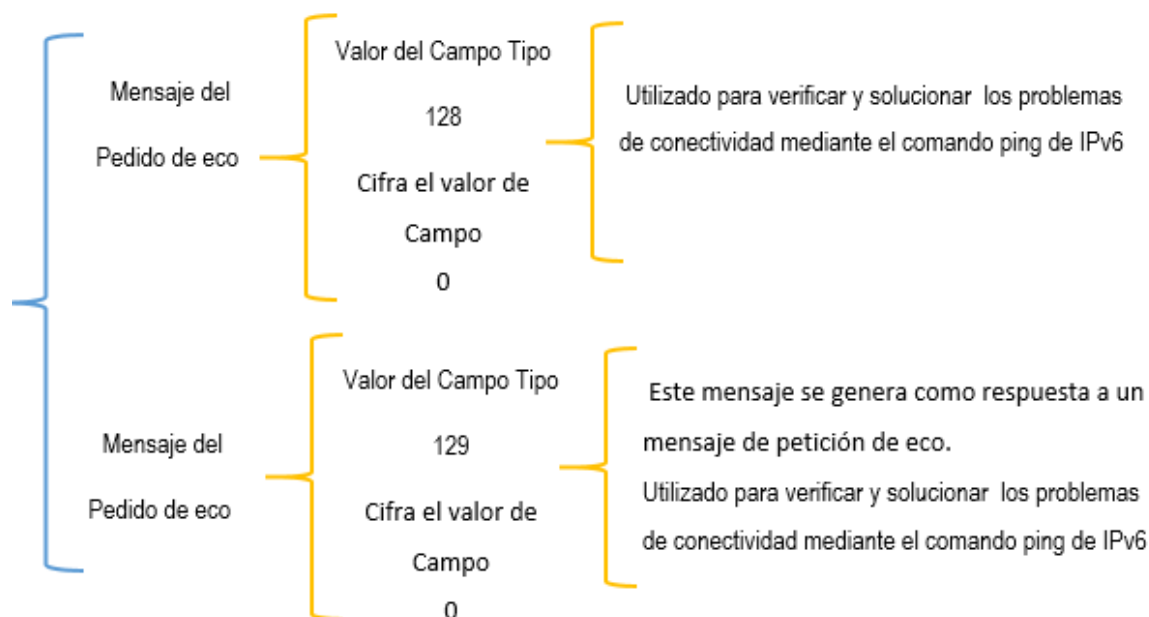


Figura 52: Mensajes de Información.



Formato:

- **Identificador y Secuencia** (16bits cada uno): Sirven para identificar las respuestas, dependen de la implementación de ping.
- **Datos:** deben copiarse en la respuesta.

Type: 128 or 129	Code: 0	Checksum
Identifier		Sequence number
Optional data Sent by the request message; repeated by the reply message		

Figura 53: Formato de mensajes de información.

3.5 Descubrimiento de Vecinos

El protocolo de descubrimiento de vecinos “Neighbor Discovery Protocol”, NDP) es un protocolo necesario para el correcto funcionamiento de las redes IPv6. Es el encargado de descubrir otros nodos en el enlace, realizar la resolución de direcciones IPv6 y direcciones MAC, encontrar los “routers” disponibles y mantener información actualizada sobre el estado de los caminos hacia otros nodos.

Este protocolo realiza funciones para IPv6 similares a las realizadas por ARP en IPv4, para el intercambio de información utiliza mensajes ICMPv6 permitiendo:

- La configuración de interfaces (autoconf.)
- La detección de duplicados (DAD-Duplicate Address Detection).
- Detección de direcciones de capa de enlaces
- Detección de vecinos inalcanzables (NUD -Neighbor Unreachability Detection).
- Redireccionamiento.

Mensajes de ICMPv6 de detección de vecino

1) Mensaje de solicitud del routers: Los hosts envían mensajes de solicitud de routers para hacer que los routers generen mensajes de anuncio de routers rápidamente.

- Valor del Campo Tipo: 133
- Cifre el valor de campo: 0

Este mensaje se genera en las siguientes situaciones:



- Tras la activación de un interfaz para detectar los encaminadores y auto configurarlo
- La dirección Multicast destino es FF02::2, todos los encaminadores

Formato:

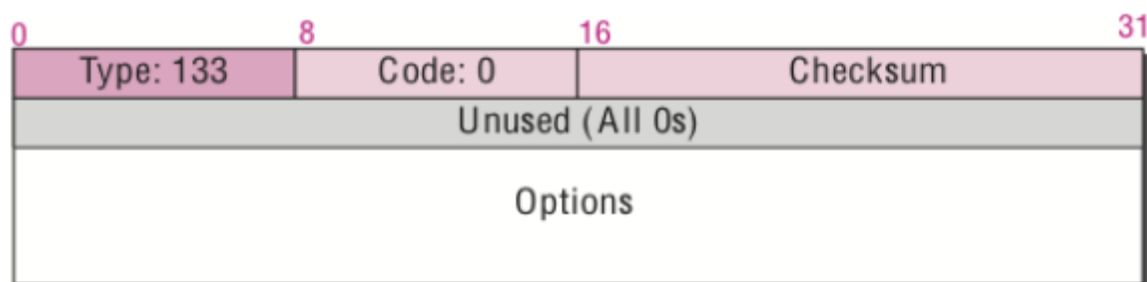


Figura 54: Mensaje de solicitud del routers.

2) Mensaje de anuncio del routers

- Los envían los routers para anunciar su presencia en la red, Estos mensajes se envían en las siguientes situaciones:
- Periódicamente, para anunciar a los hosts el router en la red. El mensaje se envía a la dirección Multicast FF02::1 (todos los hosts)
- Como respuesta a un mensaje de Router Solicitation de un host. El mensaje se envía a la dirección unicast del host solicitante.

Valor:

- Valor del Campo Tipo: 134
- Cifre el valor de campo: 0

Formato:

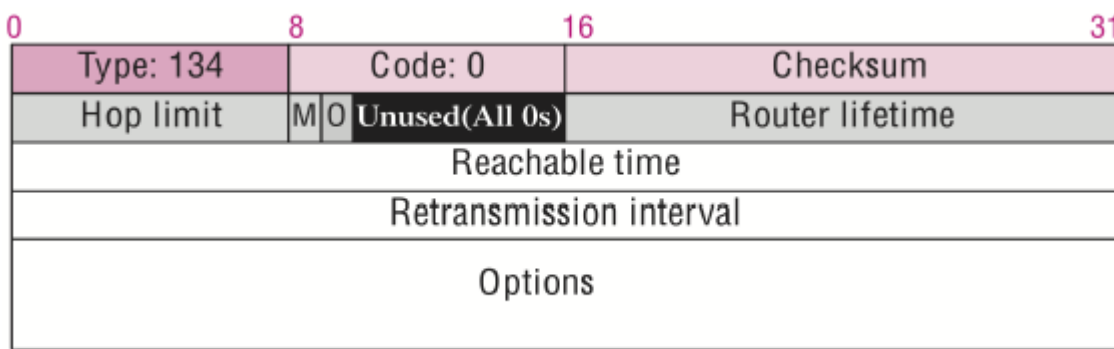


Figura 55: Mensaje de anuncio del routers

- M: Configuración vía DHCPv6
- O: Otra configuración vía DHCPv6
- Opciones: Dirección de enlace del interfaz del router, MTU y prefijo.



3) Mensaje de solicitud vecino: Los Nodos envían las solicitudes de vecino para pedir el link-layer Address de un nodo de destino mientras que también proporcionan a su propio link-layer Address a la blanco.

- Resolución de direcciones (equivalente a ARP en IPv4)
- Detección de direcciones duplicadas
- Detección de vecino alcanzable
- Mensajes ICMPv6 Neighbor Solicitation (135) y Neighbor Advertisement (136).

Valor:

- Valor del Campo Tipo: 135
- Cifre el valor de campo: 0

Este mensaje se genera en las siguientes situaciones:

- Averiguar la dirección física asociada a una dirección IP. Dirección Multicast destino es FF02::1 (todos los hosts del enlace local).
- Determinar si un nodo vecino sigue siendo alcanzable. Dirección unicast del host
- Detectar si la dirección IP está duplicada, en el proceso de autoconfiguración.

Formato:

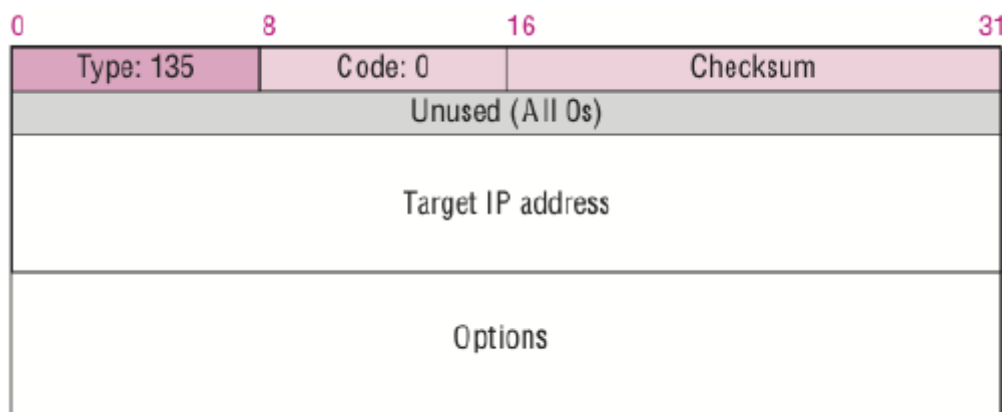


Figura 56: Mensaje de solicitud vecino

5) Mensaje de anuncio de vecino: Un nodo envía anuncios vecinos en respuesta a las solicitudes del vecino y envía anuncios no solicitados al vecino para poder propagar nueva información rápidamente (que no es confiable).

- Descubrimiento de routers
- Anuncio de prefijos y otra información de configuración de la red
- Mensajes ICMPv6 Router Solicitation (133) y Router Advertisement (134)

**Valor:**

- Valor del Campo Tipo: 136
- Cifre el valor de campo: 0

Este mensaje se genera en las siguientes situaciones:

- Respuesta a un mensaje de Neighbor Solicitation (ARP reply de IPv4). La dirección destino es la dirección unicast del destinatario.
- Cambio en la dirección física de un host, para notificar el cambio. La dirección Multicast destino

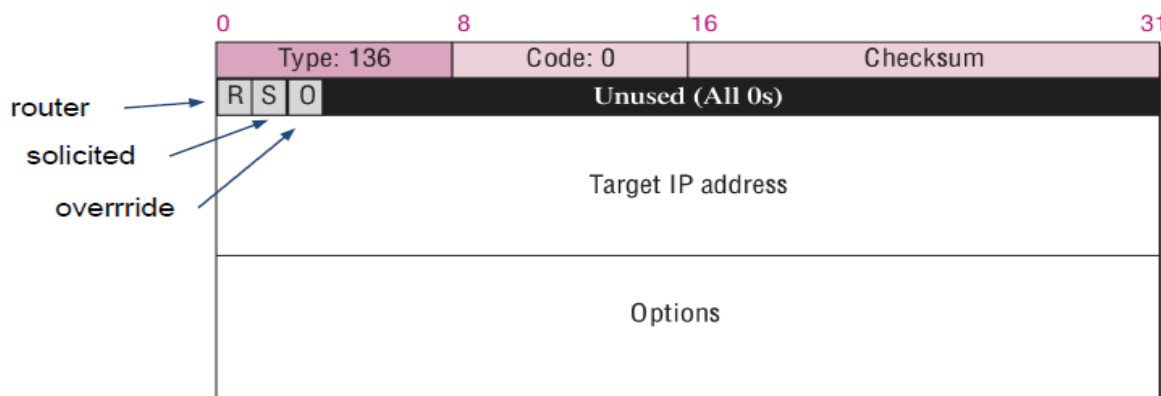


Figura 57: Mensaje de anuncio de vecino.

FF02::1

6) Redirección:

- Notificar a un host una ruta más adecuada para alcanzar un determinado destino
- Mensaje ICMPv6 Redirect.

Valor:

- Valor del Campo Tipo: 137.
- Cifre el valor de campo: 0

4. Autoconfiguración

La **autoconfiguración** es el conjunto de pasos por los cuales un host decide como auto configurar sus interfaces en **IPv6**. Este mecanismo es el que nos permite afirmar que **IPv6** es “Plug & Play”

El proceso incluye la creación de una dirección de enlace local, verificación de que no está duplicada en dicho enlace y determinación de la información que ha de ser auto configurada (direcciones y otra información).

Las direcciones pueden obtenerse de forma totalmente manual, mediante **DHCPv6** (statefull o configuración predeterminada), o de forma automática (stateless o descubrimiento automático, sin intervención).

Este protocolo define el proceso de generar una dirección de enlace local, direcciones globales y locales de sitio, mediante el procedimiento automático (stateless). También define el mecanismo para detectar direcciones duplicadas.

4.1 Autoconfiguración “stateless” (sin intervención).

No requiere ninguna configuración manual del host, configuración mínima (o ninguna) de routers, y no precisa servidores adicionales. Permite a un host generar su propia dirección mediante una combinación de información disponible localmente e información anunciada por los routers.

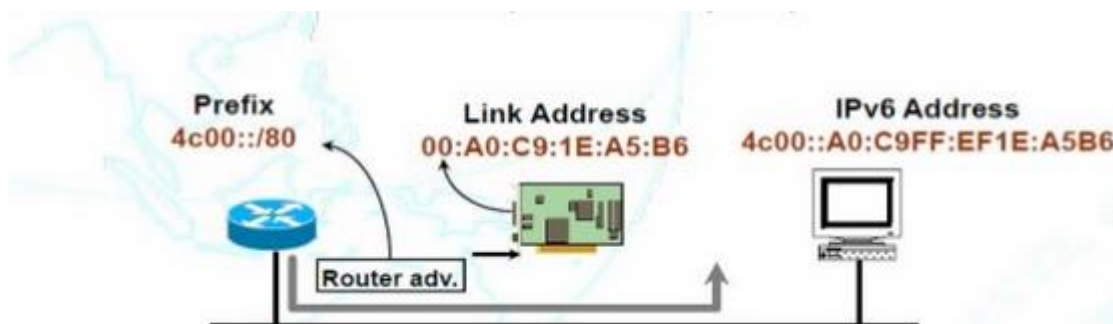


Figura 58: Autoconfiguración stateless.

Los routers anuncian los prefijos que identifican la subred (o subredes) asociadas con el enlace, mientras el host genera un “identificador de interfaz”, que identifica de forma única la interfaz en la subred. La dirección se compone por la combinación de ambos campos. En ausencia de router, el host sólo puede generar la dirección de enlace local, aunque esto es suficiente para permitir la comunicación entre nodos conectados al mismo enlace.

4.2 Autoconfiguración “stateful” (predeterminada).

El host obtiene la dirección de la interfaz y/o la información, también parámetros de configuración desde un servidor. Los servidores mantienen una base de datos con las direcciones que han sido asignadas a cada host.

Ambos tipos de **autoconfiguración** (stateless y stateful), se complementan. Un host puede usar autoconfiguración sin intervención (**stateless**), para generar su propia dirección, y obtener el resto de

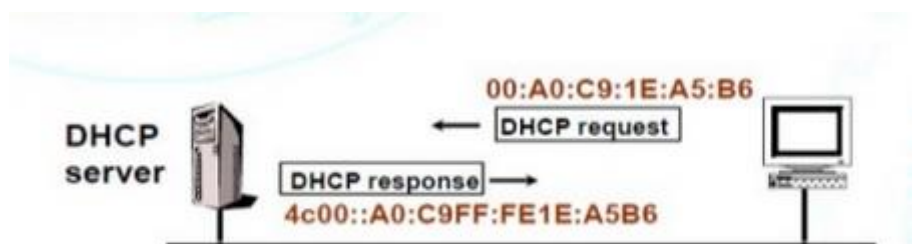


Figura 59: Autoconfiguración statefull.

parámetros mediante autoconfiguración predeterminada (**statefull**).

El mecanismo de **autoconfiguración** “sin intervención” se emplea cuando no importa la dirección exacta que se asigna a un host, sino tan sólo asegurarse que es única y correctamente enrutable.

El mecanismo de autoconfiguración predeterminada, por el contrario, nos asegura que cada host tiene una determinada dirección asignada manualmente.

Cada dirección es cedida a una interfaz durante un tiempo predefinido (posiblemente infinito). Las direcciones tienen asociado un tiempo de vida, que indican durante cuánto tiempo está vinculada dicha dirección a una determinada interfaz. Cuando el tiempo de vida expira, la vinculación se invalida y la dirección puede ser reasignada a otra interfaz en cualquier punto de Internet.

Para gestionar la expiración de los vínculos, una dirección pasa a través de dos fases diferentes mientras está asignada a una interfaz. Inicialmente una dirección es “preferred” (preferida), lo que significa que su uso es arbitrario y no está restringido. Posteriormente, la dirección es “deprecated” (desaprobada), en anticipación a que el vínculo con su interfaz actual va a ser anulado.

Para asegurarse de que todas las direcciones configuradas son únicas, en un determinado enlace, los nodos ejecutan un algoritmo de detección de direcciones duplicadas, antes de asignarlas a una interfaz. Este algoritmo es ejecutado para todas las direcciones, independientemente de que hayan sido obtenidas mediante autoconfiguración stateless o statefull.

La **autoconfiguración** está diseñada para hosts, no para routers, aunque ello no implica que parte de la configuración de los routers también pueda ser realizada automáticamente (generación de direcciones de



enlace local). Además, los routers también tienen que “aprobar” el algoritmo de detección de direcciones duplicadas.

4.3 Ejemplo de autoconfiguración.

La topología siguiente está formada por 1 router, 1 switch y 3 pc en ello se incluye 2 laptops, se utilizara un direccionamiento IPv6. La red que está presente en el router usan un prefijo /64. Se mostrará el proceso de autoconfiguración de ipv6 en las pcs que están conectadas al router.

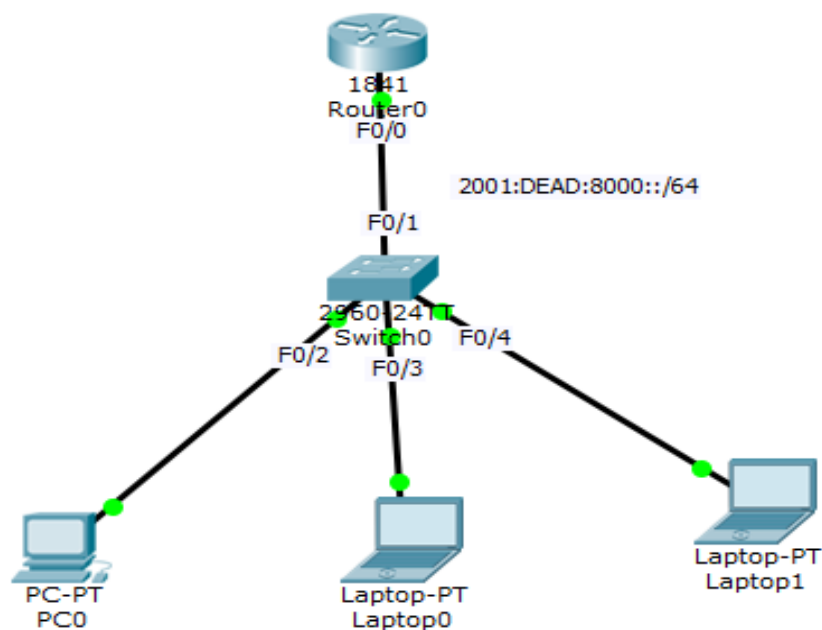


Figura 60: Topología de autoconfiguración.

Activar el ruteo por ipv6.

```
Router>en
Router#conf t
Enter configuration commands, one per line. End with CNTL/Z.
Router(config)#ipv6 unicast-routing
Router(config)#
```

Configuración de cada dispositivo dentro de la red

```
R0>en
R0#conf t
Enter configuration commands, one per line. End with CNTL/Z.
R0(config)#int f0/0
R0(config-if)#ipv6 address 2001:DEAD:8000::1/64
R0(config-if)#no shutdown
R0(config-if)#
```



Después de haber terminado con lo anterior se llevará a cabo el proceso de autoconfiguración IPv6 en las pcs y laptops respectivamente.

PC0

IPv6 Configuration	
☐ DHCP ☒ Auto Config ☐ Static	
IPv6 Address	2001:DEAD:8000:0:260:70FF:FE83:E454 / 64
Link Local Address	FE80::260:70FF:FE83:E454
IPv6 Gateway	FE80::2D0:BAFF:FE1E:5B01

Laptop0

IPv6 Configuration	
☐ DHCP ☒ Auto Config ☐ Static	
IPv6 Address	2001:DEAD:8000:0:250:FFF:FEEC:4569 / 64
Link Local Address	FE80::250:FFF:FEEC:4569
IPv6 Gateway	FE80::2D0:BAFF:FE1E:5B01

Laptop1

IPv6 Configuration	
☐ DHCP ☒ Auto Config ☐ Static	
IPv6 Address	2001:DEAD:8000:0:2D0:BAFF:FE4D:6273 / 64
Link Local Address	FE80::2D0:BAFF:FE4D:6273
IPv6 Gateway	FE80::2D0:BAFF:FE1E:5B01



5. Enrutamiento con IPv6.

5.1 Tipos de protocolos de enrutamiento

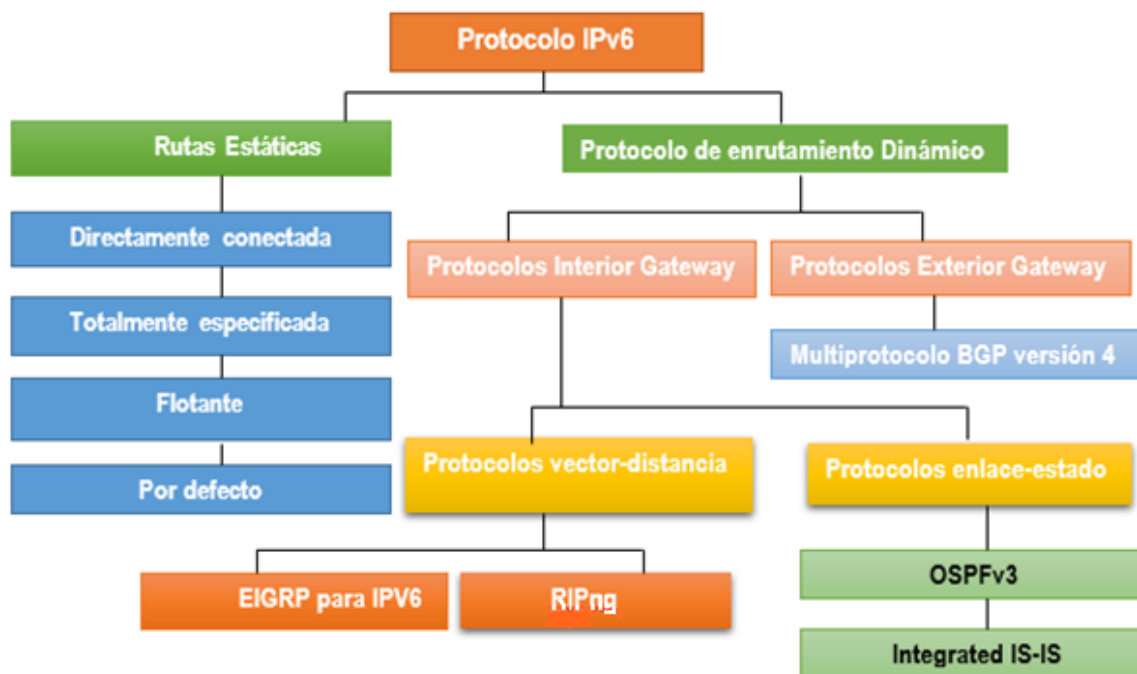


Figura 61: Tipo de enrutamiento IPv6

5.2 IGP (Protocolos de Gateway interior)

Los protocolos de Gateway interior (IGP), distribuyen la información de los routers dentro de sistemas autónomos, pueden clasificarse en dos tipos:

- Protocolos de enrutamiento vector distancia
- Protocolos de enrutamiento de Link-state (estado de enlace).

Funcionamiento del protocolo de enrutamiento vector distancia: Significa que las rutas se publican como vectores de distancia y dirección. La distancia se define en términos de una métrica como el conteo de saltos y la dirección es simplemente el router del siguiente salto o la interfaz de salida. El protocolo vector distancia generalmente usan el algoritmo Bellman-Ford para la determinación del mejor camino.

Funcionamiento del protocolo de link-state: A diferencia de la operación del protocolo de enrutamiento vector distancia, un router configurado con de link-state puede crear una "vista completa" o topología de la red al reunir información proveniente de todos los demás routers.



Con algunos protocolos de enrutamiento vector distancia, los routers envían actualizaciones periódicas de su información de enrutamiento a sus vecinos. Los protocolos de enrutamiento de link-state no usan actualizaciones periódicas. Luego de que la red ha convergido, la actualización del link-state sólo se envía cuando se produce un cambio en la topología.

Protocolos de enrutamiento vector distancia	Protocolos de enrutamiento de Link-state
➤ EGRP IPv6. ➤ RIPng.	➤ OSPFv3. ➤ Integrated IS-IS

Tabla 7: Protocolos IGP

5.3 EGP (Protocolo de pasarela externo).

En el caso de requerir una comunicación fuera del sistema autónomo EGP, es un protocolo estándar que se utiliza para el intercambio de información de enrutamiento por medio de AS's. Los Gateway en EGP sólo pueden enviar información para el acceso a redes de su sistema autónomo. El Gateway compila toda esta información, por medio de un protocolo IGP. En EGP se trabaja con BGP.

5.4 Características de los protocolos de enrutamiento IGP y EGP

Los IGP: Se usan para el enrutamiento dentro de un dominio de enrutamiento, aquellas redes bajo el control de una única organización. Un sistema autónomo está comúnmente compuesto por muchas redes individuales que pertenecen a empresas, escuelas y otras instituciones. Un IGP se usa para enrutar dentro de un sistema autónomo, y también se usa para enrutar dentro de las propias redes individuales.

Los protocolos de enrutamiento, y más específicamente el algoritmo utilizado por ese protocolo de enrutamiento, utilizan una métrica para determinar el mejor camino hacia una red. La métrica utilizada por el protocolo de enrutamiento RIP es el conteo de saltos, que es el número de routers que un paquete debe atravesar para llegar a otra red. OSPF usa el ancho de banda para determinar la ruta más corta.

Los EGP: Están diseñados para su uso entre diferentes sistemas autónomos que están controlados por distintas administraciones. El BGP es el único EGP actualmente viable y es el protocolo de enrutamiento que usa Internet. El BGP es un protocolo vector ruta que puede usar muchos atributos diferentes para medir las rutas. En el ámbito del ISP, con frecuencia hay cuestiones más importantes que la simple elección de la ruta más rápida.

5.5 Diferencias entre IGP y EGP

IGP	EGP
Descubrimiento automático de vecinos.	Vecinos son configurados específicamente.
Confianza en la información de los enrutadores que corren el IGP.	Conexión a redes externas.
Prefijos van a todos los enrutadores que corren el IGP.	Define fronteras administrativas.
Conecta enrutadores dentro de una AS.	Conecta sistemas autónomos.
Lleva solo las direcciones de infraestructura del ISP.	Lleva los prefijos de los clientes.
ISPs tratan de mantener el tamaño de las tablas del IGP bajo para eficiencia y estabilidad.	Lleva los prefijos del Internet.
	EGPs son independientes de la topología de la red del ISP.

Tabla 8: IGP vs EGP.

5.6 Sistemas Autónomos.

Conocido como dominio de enrutamiento, es un conjunto de routers que se encuentran bajo una administración común. Algunos ejemplos típicos son la red interna de una empresa y la red de un proveedor de servicios de Internet. Debido a que Internet se basa en el concepto de sistema autónomo, se requieren dos tipos de protocolos de enrutamiento: protocolos de enrutamiento interior y exterior.

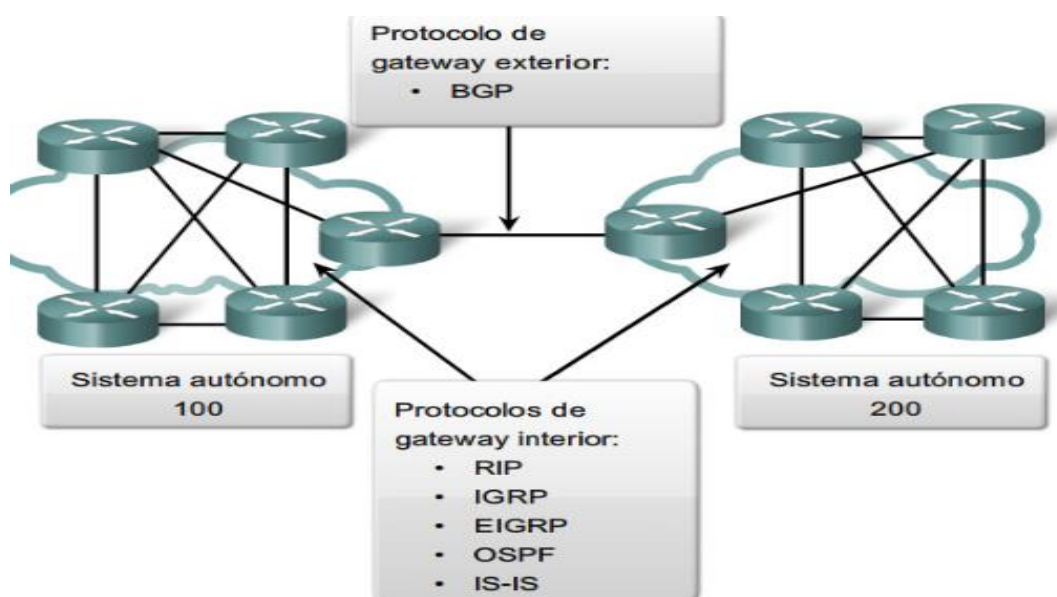


Figura 62: Comparación entre IGP y EGP

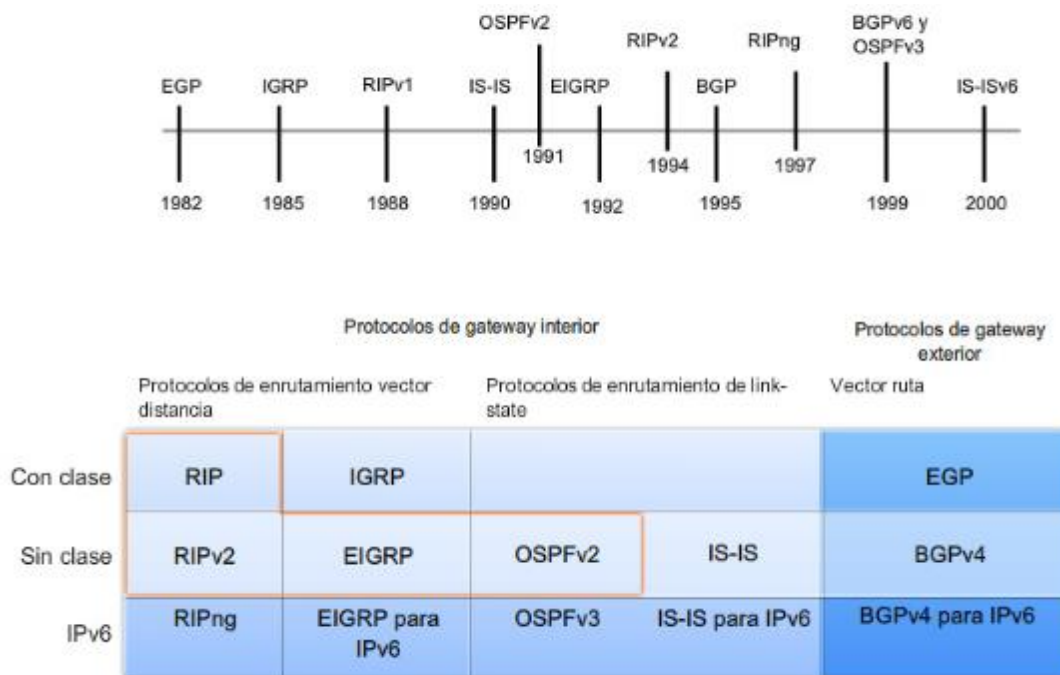


Figura 63: Evolución de los protocolos

5.7 Enrutamiento Estático.

Las rutas estáticas se definen administrativamente y establecen rutas específicas que han de seguir los paquetes para pasar de un puerto de origen hasta un puerto de destino. Se establece un control preciso del enrutamiento según los parámetros del administrador.

Las rutas estáticas se clasifican de la siguiente manera:

Ruta estática directamente conectada

Una ruta estática directamente conectada es creada cuando especificamos únicamente la interface de salida.

```
Router(config)#
```

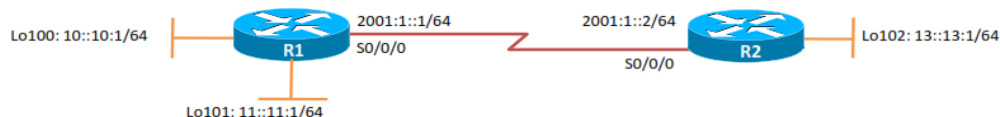
```
ipv6 route ipv6-prefix/prefix-length  
 {ipv6-address | interface-type interface-number [ipv6-address]}  
 [administrative-distance]
```

Figura 64: Configuración de ruta estática directamente conectada.

- El parámetro **ipv6-prefix/prefix-length**: Identifica la red de destino IPv6 y su longitud de prefijo.
- El parámetro **interface-type interface-number**: Especifica el interface a través del cual la red destino es alcanzada.



Ruta estática directamente conectada IPv6



```
R1# config t
R1(config)# ipv6 route 13::/64 s0/0/0
R1(config)# exit
R1# show ipv6 route static
IPv6 Routing Table - 9 entries
Codes: C - Connected, L - Local, S - Static, R - RIP, B - BGP
       U - Per-user Static route
       I1 - ISIS L1, I2 - ISIS L2, IA - ISIS interarea, IS - ISIS summary
       O - OSPF intra, OI - OSPF inter, OE1 - OSPF ext 1, OE2 - OSPF ext 2
       ON1 - OSPF NSSA ext 1, ON2 - OSPF NSSA ext 2
S      13::/64 [1/0]
       via ::, Serial0/0/0
R1#
```

Una ruta estática directamente conectada a la red 13::13:1/64 es configurada en el router R1.

Ruta estática IPv6 totalmente especificada.

Una ruta estática totalmente especificada es creada cuando especificamos: La interface de salida y la dirección IP del siguiente salto, este método evita una búsqueda recursiva.

Router(config)#

```
ipv6 route ipv6-prefix/prefix-length
{ipv6-address | interface-type interface-number [ipv6-address]}
[administrative-distance]
```

Figura 65: Ruta estática IPv6 totalmente especificada

Una ruta estática totalmente especificada a la red 13::13:1/64 es configurada en el router R1.

Ruta estática IPv6 totalmente especificada



```
R1# config t
R1(config)# ipv6 route 13::/64 s0/0/0 2001:1::2
R1(config)# exit
R1# show ipv6 route static
IPv6 Routing Table - Default - 8 entries
Codes: C - Connected, L - Local, S - Static, U - Per-user Static route
       B - BGP, M - MIPv6, R - RIP, I1 - ISIS L1
       I2 - ISIS L2, IA - ISIS interarea, IS - ISIS summary, D - EIGRP
       EX - EIGRP external
       O - OSPF Intra, OI - OSPF Inter, OE1 - OSPF ext 1, OE2 - OSPF ext 2
       ON1 - OSPF NSSA ext 1, ON2 - OSPF NSSA ext 2
S      13::/64 [1/0]
       via 2001:1::2, Serial0/0/0
R1#
```

**Observación: Ruta estática IPv6 recursiva.**

Una ruta estática recursiva es configurada cuando especificamos la dirección IP del siguiente salto del vecino, esto permite que el routers realice una segunda ruta de búsqueda para resolver la interface de salida y especificar la siguiente dirección de salto, típicamente, las rutas estáticas recursivas deben ser evitadas.

```
Router(config)#
```

```
ipv6 route ipv6-prefix/prefix-length  
{ipv6-address | interface-type interface-number [ipv6-address]}  
[administrative-distance]
```

Figura 66: Observación de ruta estática IPv6 recursiva

Ruta estática IPv6 flotante

Una ruta estática flotante es normalmente configurada cuando hay múltiples caminos a la red destino y una ruta activa de respaldo es requerida para soportar rutas descubiertas por IGP.

- Únicamente es añadida a la tabla de enrutamiento si la entrada IGP es borrada.

El parámetro administrative-distance especifica el valor de la ruta, que debe ser más alta que la ruta IGP en la tabla de enrutamiento.

- Por defecto el valor es 1, para que las rutas estáticas tenga preferencia sobre cualquier otra ruta excepto las rutas conectadas.

```
Router(config)#
```

```
ipv6 route ipv6-prefix/prefix-length  
{ipv6-address | interface-type interface-number [ipv6-address]}  
[administrative-distance]
```

Figura 67: Ruta estática IPv6 flotante.

Ejemplo:

- Por ejemplo, R1 es configurada con una ruta estática flotante especificando una distancia administrativa de 130 para la LAN de R2.



- Si un IGP ya tiene una entrada en la tabla de enrutamiento IPv6 para esta LAN, entonces la ruta estática debe únicamente aparecer en la tabla de enrutamiento si la entrada IGP fuera borrada.

Ruta estática IPv6 flotante



```
R1# config t
R1(config)# ipv6 route 13::/64 130
R1(config)# exit
R1#
```

Ruta estática IPv6 por defecto.

Las Direcciones IPv6 también tiene una ruta estática por defecto similar a la de IPv4 por defecto (0.0.0.0), en su lugar se usa la notación `::/0` para especificar todas las redes.

```
Router(config)#
```

```
ipv6 route ::/0
{ipv6-address | interface-type interface-number [ipv6-address]}
[administrative-distance]
```

Figura 68: Ruta estática IPv6 por defecto

Ruta estática IPv6 por defecto



```
R2# config t
R2(config)# ipv6 route ::/0 s0/0/0
R2(config)# exit
R2# show ipv6 route static
IPv6 Routing Table - 9 entries
Codes: C - Connected, L - Local, S - Static, R - RIP, B - BGP
       U - Per-user Static route
       I1 - ISIS L1, I2 - ISIS L2, IA - ISIS interarea, IS - ISIS summary
       O - OSPF intra, OI - OSPF inter, OE1 - OSPF ext 1, OE2 - OSPF ext 2
       ON1 - OSPF NSSA ext 1, ON2 - OSPF NSSA ext 2
S      ::/0 [1/0]
       via ::, Serial0/0/0
R2#
```

Por ejemplo, una ruta estática por defecto es especificada mediante la entrada `::/0` que es configurada en el router R2 para poder alcanzar todas las otras redes conectadas a R1.



6. DHCPv6

El Protocolo de configuración dinámica de host (DHCP) se diseñó para encargarse de la asignación de direcciones IP y otra información de red a los equipos, de forma que puedan comunicarse en la red automáticamente. DHCPv6, puede proporcionar configuración de direcciones con estado o configuración sin estado a hosts de IPv6. Los hosts de IPv6 pueden utilizar varios métodos de configuración de direcciones.

En comparación con otros métodos de asignación de direcciones IPv6 (como la configuración manual y la configuración automática de direcciones sin estado). DHCPv6 puede realizar lo siguiente:

- Registro de direcciones asignadas a los hosts y asignar direcciones específicas a los hosts, lo que facilita la gestión de redes.
- Asignar prefijos a los dispositivos, facilitando así la configuración automática y la gestión de toda la red.
- Asignar otros parámetros de configuración, como las direcciones de servidor DNS y nombres de dominio, a los ejércitos.

DHCPv6 funciona sobre el protocolo de transporte UDP. El cliente utiliza una dirección link-local u otra determinada a través de otros mecanismos para transmitir y recibir los mensajes DHCPv6, puede trabajar de forma conjunta con el mecanismo “stateless”. El administrador de red determina que procesos se van a emplear a través de los mensajes “RA” de ICMPv6.

Los servidores DHCPv6 reciben mensajes de los clientes utilizando una dirección reservada Multicast. Un cliente DHCPv6 transmite la mayoría de los mensajes a la dirección anteriormente mencionada por lo que no es necesario que el cliente sea configurado con dirección unicast del servidor DHCPv6.

6.1 Arquitectura Cliente-servidor.

La arquitectura cliente-servidor es utilizada como base del funcionamiento de ese protocolo. En cada red debe haber un servidor capaz de decidir sobre la configuración de cada una de las interfaces de red.

Todos los mensajes DHCPv6 enviados entre clientes y servidor comparten un formato de cabecera fijo e idéntico y un formato variable en el área de opciones. Las opciones se guardan seguidamente sin relleno entre ellas. En la siguiente figura se puede ver un diagrama de la composición de dicho mensaje en el que se pueden identificar los siguientes campos:

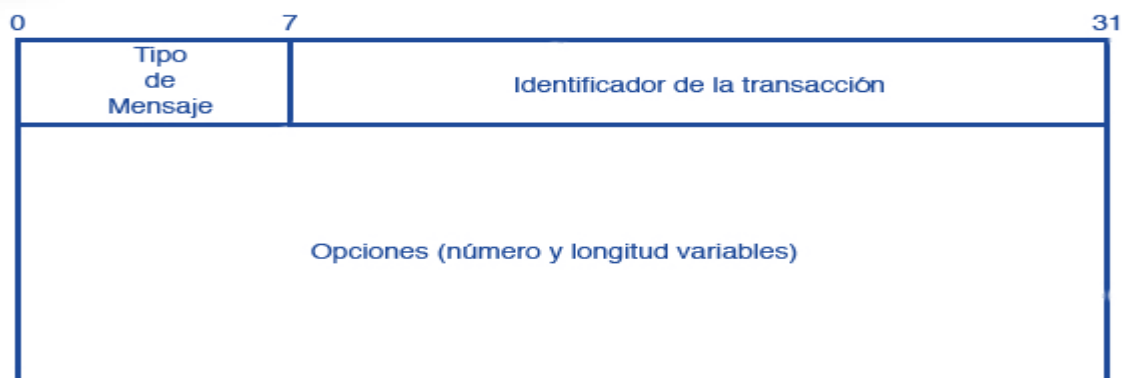


Tabla 9: Formato de mensaje cliente-servidor.

- Tipo de mensaje. Identifica el tipo de mensaje DHCP.
- Identificador de transacción. Número identificativo de un intercambio de mensajes.
- Opciones. Este campo es variable.

6.2 Identificador Único DHCP (DUID).

Cada cliente o servidor DHCPv6 tiene un DUID. Los servidores DHCPv6 usan DUIDs para identificar a los clientes para la selección de parámetros de configuración y para asociar las IAs con los clientes. Los clientes usan DUIDs para identificar al servidor en mensajes en los que el servidor debe ser identificado.

Tanto cliente como servidor tratan los DUIDs como valores opacos y sólo deben compararlos por igualdad. Ambos en ningún caso interpretan los DUIDs. El DUID es transportado en el campo opciones ya que tiene una longitud variable y no se requiere en todos los mensajes DHCPv6.

El DUID consiste en un campo de dos octetos representando el tipo de código seguido por un número variables de octetos que representan el identificador real. Un DUID no puede ser más largo de 128 bytes (sin incluir el tipo de código).

Formato de DUID

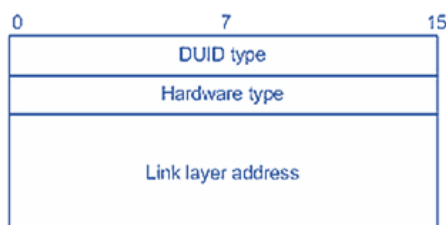


Tabla 10: Formato de DUID-LL



Actualmente, se utiliza un DUID basado en la dirección de capa de enlace (DUID-LL) definido en RFC 3315 para identificar un dispositivo DHCPv6. La figura muestra el formato DUID-LL, donde:

- **DUID type:** El dispositivo es compatible con DUID-LL como el tipo DUID con el valor de 0x0003.
- **Hardware type:** El dispositivo es compatible con Ethernet como el tipo de hardware con el valor de 0x0001.
- **Link layer address:** Su valor es la dirección MAC del puente del dispositivo.

6.3 Mensajes DHCPv6.

Al igual que en DHCP para IPv4, DHCPv6 utiliza mensajes de Protocolo de datagramas de usuario (UDP). Los clientes DHCPv6 escuchan mensajes DHCP en el puerto 546 de UDP. Los agentes de retransmisión y los servidores DHCPv6 escuchan mensajes en el puerto 547 de UDP. La estructura de mensajes DHCPv6 es mucho más sencilla que la de DHCP para IPv4, que tuvo sus orígenes en el protocolo BOOTP para ofrecer compatibilidad con estaciones de trabajo sin disco.

La siguiente tabla muestra los tipos de mensaje DHCPv6 así como su función:

Mensaje	Msg-type	Descripción
SOLICIT	1	Un cliente manda un mensaje SOLICIT para localizar servidores.
ADVERTISE	2	El servidor en un mensaje ADVERTISE para indicar que está disponible para el servicio, en respuesta al mensaje de solicitud recibido de un cliente.
REQUEST	3	Un cliente manda un mensaje REQUEST para solicitar los parámetros de configuración, incluyendo la dirección IP, de un servidor específico.
CONFIRM	4	Un cliente envía un mensaje CONFIRM a cualquier servidor disponible para determinar si las direcciones que se asignaron siguen siendo válidas.
RENEW	5	Un cliente manda un mensaje RENEW al servidor que originalmente proporcionó la dirección del cliente y los parámetros de configuración para



		extender los tiempos de vida en la dirección asignada y actualizar los parámetros de configuración.
REBIND	6	Un cliente envía el mensaje REBIND a cualquier servidor disponible para extender los tiempos de vida en la dirección asignada a los clientes y para actualizar otros parámetros de configuración. Este mensaje solo se envía cuando el cliente no obtiene respuesta del mensaje RENEW.
REPLY	7	Un servidor manda un mensaje de REPLY conteniendo las direcciones asignadas y los parámetros de configuración en respuesta a un mensaje SOLICIT, REQUEST, RENEW, REBIND recibido del cliente. También se manda este mensaje pero conteniendo únicamente los parámetros de configuración en respuesta a un mensaje INFORMATION-REQUEST. También se manda un mensaje REPLY en respuesta a un mensaje CONFIRM confirmando o denegando la validez de la dirección asignada al cliente que envía el mensaje CONFIRM. Y por último también se envía un mensaje de REPLY para realizar un asentimiento de recepción de los mensajes RELEASE y DECLINE.
RELEASE	8	Un cliente manda un mensaje de RELEASE al servidor que le asignó las direcciones IP para indicar que no usará más una o varias de las direcciones asignadas.
DECLINE	9	Un cliente envía un mensaje DECLINE para indicar al cliente que una o varias de las direcciones que tiene asignadas por el servidor están siendo utilizadas en el enlace en el que el cliente está conectado.
RECONFIGURE	10	Un servidor envía un mensaje de RECONFIGURE a un cliente para informarle de que el servidor tiene unos parámetros de configuración nuevos o actualizados, entonces el cliente iniciará una transacción RENEW/REPLY o INFORMATION-REQUEST/REPLY con el servidor para obtener la nueva configuración.
INFORMATION-REQUEST	11	Un cliente envía un mensaje INFORMATION-REQUEST al servidor para pedir los parámetros de configuración sin pedir ninguna dirección IP.
RELAY-FORW	12	Un agente de retransmisión envía un mensaje RELAY-FORWARDING al servidor, ya sea él mismo o a través de otro agente de transmisión. El mensaje



		recibido, a través de un cliente o a través de otro agente de retransmisión se encapsula en una opción del mensaje RELAY-FORWARDING.
RELAY-REPL	13	Un servidor envía un mensaje RELAY-REPLY a un agente de retransmisión, conteniendo el mensaje que el agente de retransmisión debe entregar al cliente. Este mensaje puede ser reenviado por varios agentes de retransmisión. El servidor encapsula el mensaje al cliente como una opción dentro del mensaje RELAY-REPLY.

Tabla 11: Mensajes de DHCPv6.

Un proceso de asignación de direcciones DHCPv6 / prefijo implica dos o cuatro mensajes. La siguiente descripción muestra los procesos detallados.

➤ **Asignación rápida, participan dos mensajes:**

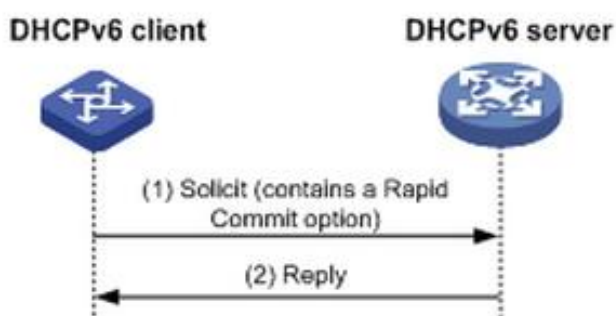


Figura 69: DHCPv6, asignación rápida con dos mensajes.

Como se muestra en la Figura el proceso de asignación rápida que implica dos mensajes es el siguiente:

- 1) El cliente DHCPv6 envía un mensaje Solicit que contiene una opción rápida Commit, solicitando que se debe preferir la asignación rápida de dirección / prefijo y otros parámetros de configuración.
- 2) Si el servidor DHCPv6 soporta asignación rápida, responde con un mensaje Reply que contiene la dirección IPv6 asignada / prefijo y otros parámetros de configuración. Si el servidor DHCPv6 no permite la asignación rápida, entonces viene la asignación en la que se implementan cuatro mensajes.

➤ **Asignación con implementación de cuatro mensajes:**

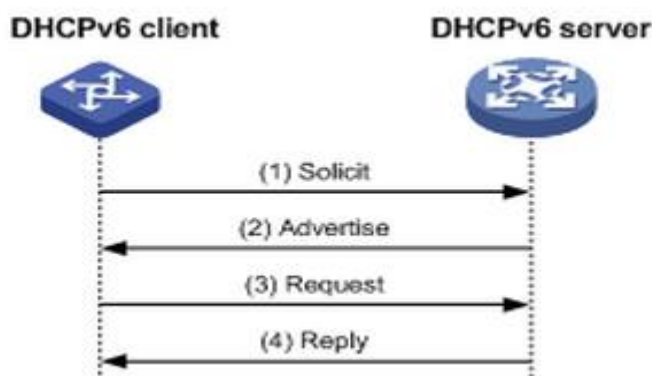


Figura 70: DHCPv6, asignación con 4 mensajes.

El proceso de asignación de la participación de cuatro mensajes es el siguiente:

- 1) El cliente DHCPv6 envía un mensaje Solicit, solicita una dirección IPv6 / prefijo y otros parámetros de configuración.
- 2) Si el servidor DHCPv6 no permite la asignación rápida, el servidor DHCPv6 responde con un mensaje de Advertise, informando al cliente DHCPv6 de la dirección / prefijo asignable y otros parámetros de configuración.
- 3) El cliente DHCPv6 puede recibir múltiples mensajes de Advertise (publicidad) ofrecidos por diferentes servidores DHCPv6, se selecciona una de las ofertas de acuerdo a la secuencia y el servidor de recepción de prioridad, y envía un mensaje Request al servidor seleccionado para la confirmación de la asignación.
- 4) El servidor DHCPv6 envía un mensaje Reply al cliente, lo que confirma que la dirección / prefijo y otros parámetros de configuración están asignados al cliente.

➤ **Dirección / prefijo, mensaje de renovación (Renew).**

La dirección / prefijo de IPv6 asignado por el servidor DHCPv6 tiene un tiempo de concesión, que depende de la duración válida. Cuando el tiempo de vida válido de la dirección / prefijo IPv6 expira, el cliente DHCPv6 no puede utilizar la dirección / prefijo IPv6 por más tiempo. Para utilizar la dirección / prefijo IPv6 más tiempo, el cliente DHCPv6 tiene que renovar el tiempo de concesión.

Usando el mensaje Renew para la dirección / prefijo.

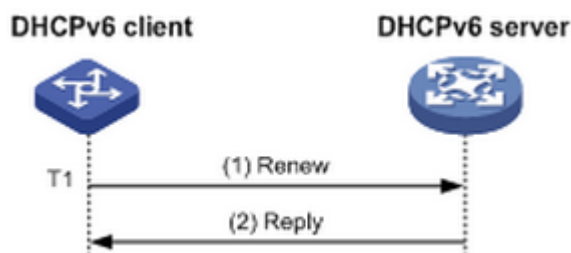


Figura 71: DHCPv6, mensaje Renew.

Como se muestra en la Figura, en T1, el cliente DHCPv6 difunde un mensaje de Renew al servidor DHCPv6 que asigna la dirección / prefijo IPv6 al cliente DHCPv6. El valor recomendado de T1 es la mitad de la vida útil preferente. A continuación, el servidor DHCPv6 responde con un mensaje de respuesta, informando si hay o no renovación.

Usando el mensaje Rebind en la renovación de la dirección / prefijo.

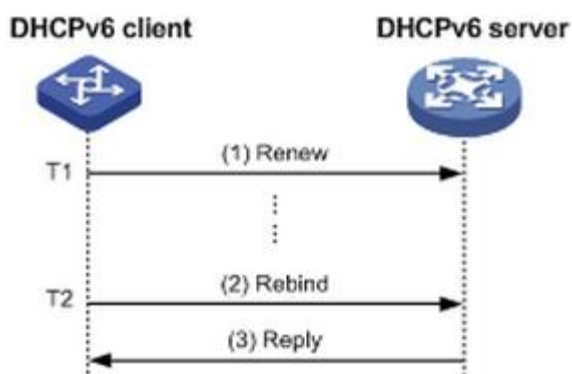


Figura 72: DHCPv6, mensaje Rebind.

Como se muestra en la Figura, si el cliente DHCPv6 no recibe respuesta desde el servidor DHCPv6 después de enviar un mensaje de Renew en T1, se difunde un mensaje Rebind a todos los servidores DHCPv6 en T2 (es decir, cuando el 80% de por vida preferido expira). A continuación, el servidor DHCPv6 responde con un mensaje de respuesta, informando si hubo renovación o no.

Si el cliente DHCPv6 no recibe ninguna respuesta de los servidores DHCPv6, el cliente deja de usar la dirección / prefijo cuando la duración válida expire.

Direcciones Multicast.

- **All_DHCP_Relay_Agents_and_Servers (FF02::1:2).** Una dirección link-local Multicast usada por el cliente para comunicarse con los agentes de transmisión y servidores vecinos. Todos los servidores y agentes de retransmisión son miembros de este grupo Multicast.
- **All_DHCP_servers (FF05::1:3).** Una dirección Site-local Multicast, usada por los agentes de retransmisión para comunicarse con los servidores, ya sea porque quiere enviar el mensaje a todos los servidores o porque no conoce la dirección unicast de los servidores.

6.4 Agente DHCPv6 relay

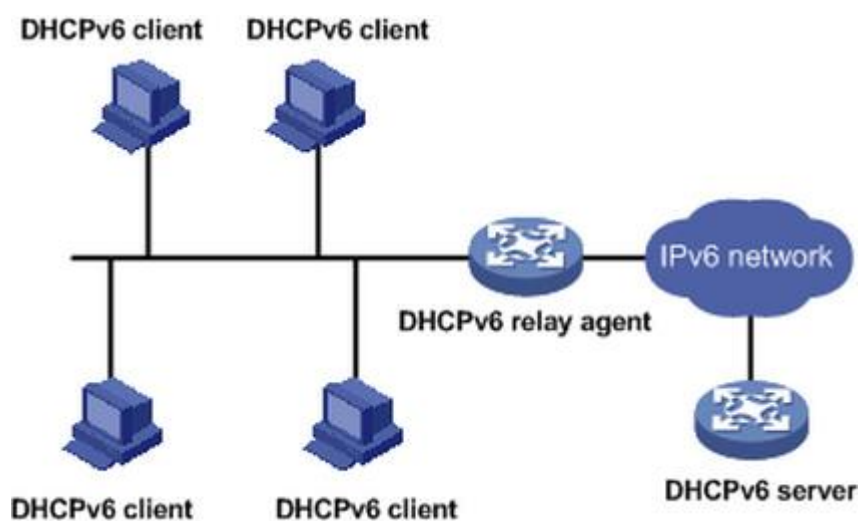


Figura 73: Aplicación típica de DHCPv6 Relay Agent.

Un cliente DHCPv6 generalmente utiliza una dirección de multidifusión para comunicarse con el servidor DHCPv6 en el enlace local para obtener una dirección IPv6 y otros parámetros de configuración. Como se muestra en la figura, si el servidor DHCPv6 reside en otra subred, el cliente DHCPv6 puede contactar con el servidor a través de un agente de retransmisión DHCPv6. Por lo tanto, no es necesario para implementar un servidor DHCPv6 en cada subred.

Operación del agente DHCPv6 Relay

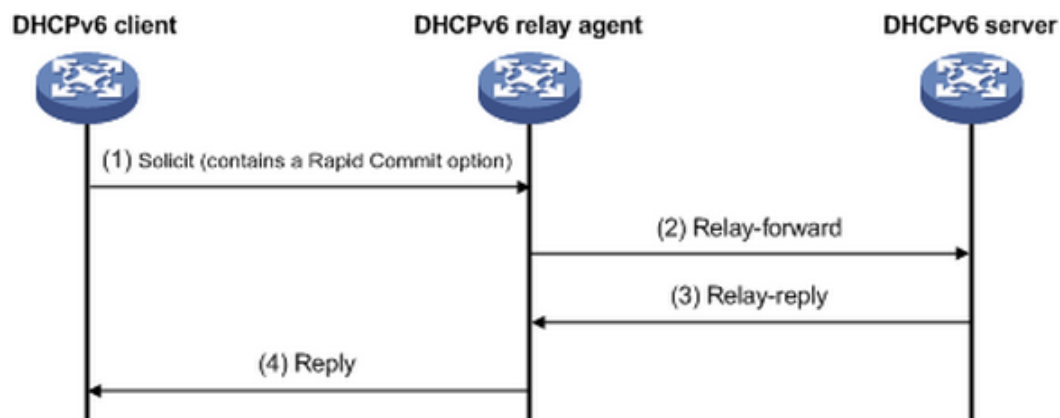


Figura 74: Proceso de funcionamiento de un DHCPv6 relay agent.

Tome el proceso de asignación rápida que implica dos mensajes como un ejemplo. La figura muestra cómo el cliente DHCPv6 obtiene la dirección IPv6 y otros parámetros de configuración de red del servidor DHCPv6 a través del agente de retransmisión DHCPv6.

1. El cliente DHCPv6 envía un mensaje Solicit que contiene la opción Rapid Commit para la dirección multicast FF02::1:2 de todos los servidores DHCPv6 y agentes de retransmisión.
2. Después de recibir el mensaje de solicitud, el agente de retransmisión DHCPv6 encapsula el mensaje en la opción de mensaje de retransmisión de un mensaje Relay-forward, y envía el mensaje al servidor DHCPv6.
3. Después de obtener el mensaje Solicit del mensaje Relay-forward, el servidor DHCPv6 selecciona una dirección IPv6 y otros parámetros necesarios, y los agrega a la respuesta que se encapsula dentro de la opción de mensajes de retransmisión de un mensaje Relay-reply. El servidor DHCPv6 envía el mensaje Relay-reply al agente de retransmisión DHCPv6.
4. El agente de retransmisión DHCPv6 obtiene la respuesta del mensaje Relay-reply y envía la respuesta al cliente DHCPv6.

Entonces el cliente DHCPv6 utiliza la dirección IPv6 y otros parámetros de red asignados por el servidor DHCPv6 para realizar la configuración de red.

Formato de mensajes entre agentes de retransmisión y servidores

Los agentes de retransmisión intercambian mensajes con servidores para retransmitir mensajes entre servidores y clientes que no están conectados en el mismo enlace. Hay dos tipos de mensajes de agente de retransmisión, pero ambos comparten el mismo formato:

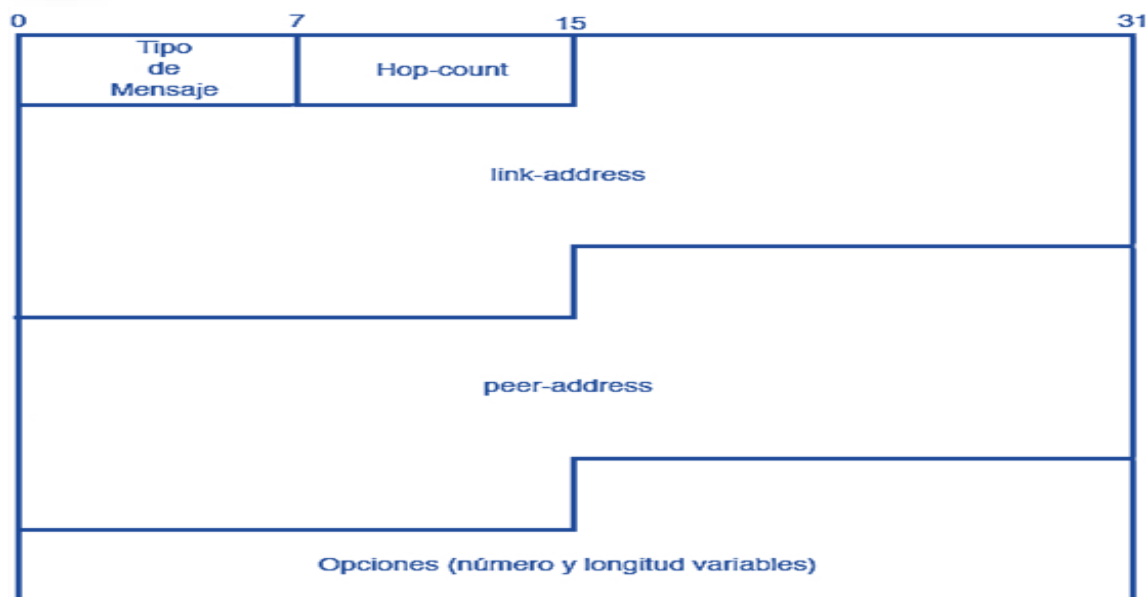


Tabla 12: Formato de mensaje entre agentes de retransmisión DHCPv6.

- Tipo de mensaje. RELAY-FORW o RELAY-REPL.
- Link-Address. Una dirección global o Site-local que será usada por el servidor para identificar el enlace en el que se encuentra el cliente.
- Peer-Address. La dirección del cliente o del agente de retransmisión desde el que se recibió el mensaje que tiene que ser retransmitido.
- Opciones. Debe incluir una opción llamada "Relay Message Options", y además puede incluir otras opciones.

6.5 Autenticación de mensajes DHCPv6

La autenticación de mensajes DHCPv6 se consigue mediante el uso de la opción Authentication. La información transportada en esta opción se puede utilizar para identificar de forma fiable el origen de un mensaje DHCPv6 y para confirmar que el contenido del mensaje DHCPv6 no ha sido manipulado. Cualquier mensaje DHCPv6 no debe incluir más de una opción de autenticación.

Seguridad de mensajes enviados entre Servidores y Agentes de Retransmisión: Los agentes de retransmisión y los servidores que intercambian mensajes de manera segura utilizan los mecanismos de IPsec para IPv6. Si un mensaje del cliente es retransmitido por varios agentes de retransmisión cada uno de los agentes debe tener relaciones de confianza establecidas por pares independientes.

Detección de Repeticiones: El campo Método de Detección de Repeticiones (RDM) determina el tipo de detección de repetición usado en el campo Replay Detection. Usar un valor cambiante, como por ejemplo la hora exacta puede reducir el peligro de los ataques de repetición.

Protocolo de Autenticación Retardada: Si el valor del campo de protocolo es 2 de la opción Authenticate el mensaje está usando el mecanismo “Delayed Authentication”. En este mecanismo el cliente pide autenticación en su mensaje SOLICIT, y el servidor responde con un mensaje ADVERTISE que contiene información de autenticación. Esta información de autenticación contiene un reto generado por el origen como un código de autenticación de mensaje (MAC) para proporcionar autenticación de mensaje y autenticación de la entidad.

El emisor del mensaje calcula el MAC utilizando el algoritmo de generación HMAC y la función de hash MD5. El mensaje DCPHv6 entero (poniendo el campo MAC en la opción de autenticación a 0) incluyendo la cabecera y el campo de opciones se usa como entrada de la función HMAC-MD5.

El receptor del mensaje calcula el HMAC-MD5 como lo hizo el emisor del mensaje y si concuerdan valida el mensaje y en caso contrario, lo descarta.

Cada cliente DHCPv6 tiene un conjunto de claves. Cada clave está identificada por <DHCP realm, DUID cliente, key-id>. Cada clave tiene un tiempo de vida limitado, por lo que no debe ser usada una vez haya pasado este tiempo.

6.6 DHCPv6 con estado y sin estado.

Configuración automática de direcciones sin estado.

Se utiliza para configurar las direcciones locales de vínculos y las direcciones no locales de vínculos adicionales mediante el intercambio de mensajes de solicitud de enrutador y anuncio de enrutador con los enrutadores vecinos.



Figura 75: Configuración automática sin estado (I)

Arranque: Construcción de Dirección de ámbito local (Link Local).

- **Direcciones Locales:** Comunicación dentro de subred (No se encaminan) y son muy útiles en redes sin router.
- **Direcciones Globales:** Se encargan de envío periódico de paquetes (Router Advertisement o RA).



Figura 76: Configuración automática sin estado (II).

Configuración automática de direcciones con estado.

Se utiliza para configurar direcciones no locales de vínculos a través del uso de un protocolo de configuración como DHCP.

Un host de IPv6 realiza la configuración de direcciones sin estado automáticamente y utiliza un protocolo de configuración (como DHCPv6) en el mensaje de anuncio de enrutador basado en las siguientes marcas y enviado por un enrutador vecino.

Al igual que en DHCP para IPv4, los componentes de una infraestructura DHCPv6 están formados por clientes DHCPv6 que solicitan configuración, servidores DHCPv6 que ofrecen configuración y agentes de retransmisión DHCPv6 que transmiten mensajes entre clientes y servidores cuando los clientes se encuentran en subredes que no tienen un servidor DHCPv6.

6.7 Ejemplo de configuración de DHCPv6

Diagrama de la topología

El enrutador DHCPv6 tiene dos interfaces, el que está conectado a R1 se utilizará para DHCPv6 con stateful y el R2 stateless.

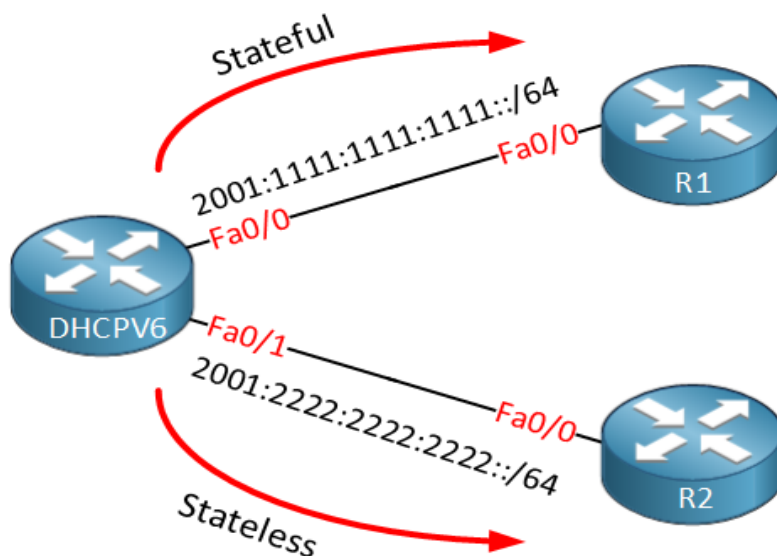


Figura 77: Topología de DHCPv6

Antes de configurar IPv6, asegúrese de que el enrutamiento de unidifusión está habilitada:

```
DHCPV6 (config) # ipv6 unicast-routing
```

Configuración de stateful DHCPv6

```
DHCPV6(config)#ipv6 dhcp pool STATEFUL
DHCPV6(config-dhcpv6)#address prefix 2001:1111:1111:1111::/64
DHCPV6(config-dhcpv6)#dns-server 2001:4860:4860::8888
DHCPV6(config-dhcpv6)#domain-name NETWORKLESSONS.LOCAL
```

El pool se llama "stateful" y además el prefijo que configura un servidor DNS (que es Google DNS) y un nombre de dominio. Para activar esto tenemos que hacer algunos cambios en la interfaz:

En la interfaz hay que añadir el comando **ipv6 dhcp server** y decirle que pool tiene que utilizar. El comando **ipv6 nd managed-config-flag** establece un indicador en el anuncio del router que dice a los anfitriones que podrían utilizar DHCPv6. El último comando que termina con **no-autoconfig** dice a los anfitriones a no utilizar la configuración sin estado.

```
DHCPV6(config)#interface FastEthernet 0/0
DHCPV6(config-if)#ipv6 address 2001:1111:1111:1111::1/64
DHCPV6(config-if)#ipv6 dhcp server STATEFUL
DHCPV6(config-if)#ipv6 nd managed-config-flag
DHCPV6(config-if)#ipv6 nd prefix 2001:1111:1111:1111::/64 14400 14400 no-
autoconfig
```



Eso es todo lo que tenemos que hacer en el servidor DHCPv6, vamos a pasar a la configuración sin estado. En primer lugar vamos a crear un pool:

```
DHCPV6(config)#ipv6 dhcp pool STATELESS  
DHCPV6(config-dhcpv6)#dns-server 2001:4860:4860::8888  
DHCPV6(config-dhcpv6)#domain-name NETWORKLESSONS.LOCAL
```

Como se puede ver que no se configuro un prefijo, habilitar en la interfaz:

```
DHCPV6(config)#interface FastEthernet 0/1  
DHCPV6(config-if)#ipv6 address 2001:2222:2222:2222::2/64  
DHCPV6(config-if)#ipv6 dhcp server STATELESS  
DHCPV6(config-if)#ipv6 nd other-config-flag
```

Utilizamos el mismo comando para activar el pool en la interfaz pero hay un elemento adicional. El **ipv6 nd other-config-flag** se requiere, ya que informar a los clientes a través de la RA (Router Advertisement) de los mensajes que tienen que utilizar DHCPv6 para recibir información adicional como el nombre de dominio y el Servidor DNS después de que utilizan la configuración automática.

Para ver las direcciones de DHCP utilizamos el siguiente comando:

```
# show ipv6 dhcp pool
```

Configuración del cliente DHCPv6

R1 será el cliente con estado y R2 es el cliente sin estado, vamos a configurar R1.

Configuración del cliente DHCPv6 con estado

Hay dos cosas que tenemos que hacer, primero necesita habilitar IPv6 en la interfaz y en segundo lugar, decir que para obtener una dirección IPv6 a través de DHCP:

```
R1(config)#interface FastEthernet 0/0  
R1(config-if)#ipv6 enable  
R1(config-if)#ipv6 address dhcp
```

Vamos a ver si tiene una dirección IPv6:



```
R1#show ipv6 interface brief
FastEthernet0/0          [up/up]
FE80::21D:A1FF:FE8B:36D0
2001:1111:1111:1111:255A:E159:32AF:5E42
```

El siguiente comando nos muestra qué más hemos recibido:

```
R1#show ipv6 dhcp interface FastEthernet 0/0
FastEthernet0/0 is in client mode
Prefix State is IDLE
Address State is OPEN
Renew for address will be sent in 11:59:10
List of known servers:
  Reachable via address: FE80::216:C7FF:FEBE:EC8
  DUID: 000300010016C7BE0EC8
  Preference: 0
Configuration parameters:
  IA NA: IA ID 0x00030001, T1 43200, T2 69120
    Address: 2001:1111:1111:1111:255A:E159:32AF:5E42/128
             preferred lifetime 86400, valid lifetime 172800
             expires at Jul 19 2014 08:30 PM (172750 seconds)
    DNS server: 2001:4860:4860::8888
    Domain name: NETWORKLESSONS.LOCAL
  Information refresh time: 0
Prefix Rapid-Commit: disabled
Address Rapid-Commit: disabled
```

El **show ipv6 dhcp interface command** nos muestra lo absoluto de DNS y la información de dominio que recibimos, esto se ve bien. Mientras tanto se puede ver esto en el servidor:



7. RIPng (RIP new Generation)

Es un protocolo de puerta de enlace interna o IGP (Interior Gateway Protocol) utilizado por los routers (encaminadores) para intercambiar información acerca de redes IP a las que se encuentran conectados. Su algoritmo de encaminamiento está basado en el vector de distancia, ya que calcula la métrica o ruta más corta posible hasta el destino a partir del número de "saltos" o equipos intermedios que los paquetes IP deben atravesar. El límite máximo de saltos en RIP es de 15, de forma que al llegar a 16 se considera una ruta como inalcanzable o no deseable.

RIPng es el protocolo que permite a los routers intercambiar información para rutas a través de una red basada en IPv6.

7.1 Generalidades de RIPng

- Es la próxima generación de RIP que adhiere soporte a IPv6.
- Se mantiene clasificado como un protocolo de enrutamiento vector distancia y usa saltos tal como RIPv1 y RIPv2.
- No está diseñado para ser utilizado en redes de gran tamaño, pero funciona bien en la mayoría de redes de tamaño pequeño y mediano.
- No soporta más de 15 saltos así como en RIPv1 y RIPv2.
- No hace autenticación de paquetes (no lo hace debido a que hace uso de IPsec).
- No usa máscaras de subred, pero utiliza una longitud de prefijo en su lugar.
- Usa direcciones Multicast FF02::9.
- No usa el puerto UDP 520 así como RIPv1 y RIPv2, pero usa el puerto UDP 521 en su lugar.

Cada router que implementa RIPng tiene una tabla de ruteo con una entrada para cada destino IPv6 alcanzable. Cada entrada debe contener al menos la siguiente información:

- El prefijo IPv6 del destino.
- Una medida que indica el costo total de obtener un datagrama desde el Router hasta el destino.
- La dirección IPv6 del próximo router en el camino al destino, llamado el próximo salto (hop).
- Un Cambio de Estado (Change Flag) que indica si la información acerca de esa ruta ha sido modificada recientemente.
- Varios relojes (timers), como un reloj de 30 segundos que apunte la transmisión de la información de la tabla de ruteo a los routers vecinos.



7.2 Formato de mensaje RIPng

RIP es un protocolo basado en UDP. Cada Router que usa RIP tiene un proceso de ruteo que envía y recibe datagramas sobre el número de puerto 521 UDP, el puerto RIP. Todas las comunicaciones entendidas para otros procesos router RIP son enviadas al puerto RIP. Todos los mensajes de actualización son enviados desde el puerto RIP. Los mensajes de actualización de ruteo no solicitados tienen igual puerto de fuente y destino que el puerto RIP. Estos envían en respuesta a un requerimiento enviado al puerto desde el cual el requerimiento vino. Las colas específicas pueden ser enviadas desde otros puertos diferentes al RIP, pero deben ser direccionados al puerto RIP sobre la máquina.

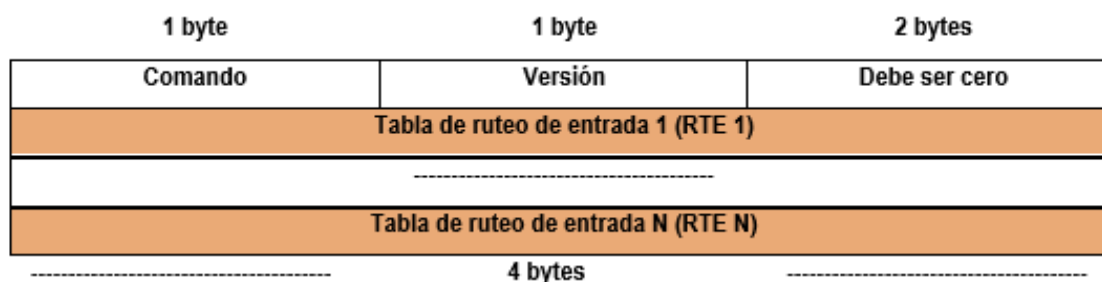


Figura 78: Formato de mensajes RIPng

Donde cada Tabla de Entrada de Ruteo (RTE) tiene el siguiente formato:

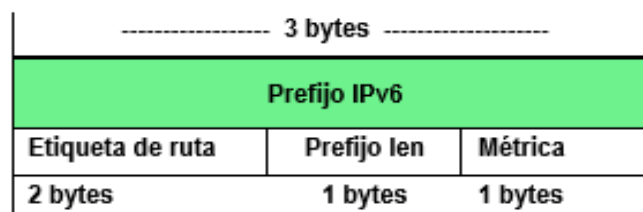


Figura 79: Ruteo RTE

- **Prefijo** = prefijo IPv6 (128 bits).
- **Etiqueta de ruta (route tag)** = atributo asignado a la ruta que debe preservarse y redistribuirse con la ruta. Permite separar las rutas RIPng internas de las que son importadas de un EGP u otro IGP.
- **Largo del prefijo** = 0 a 128.
- **Métrica** = 1 a 16 (infinito, no alcanzable).



Próximo salto RTE

RIPng también provee la habilidad de especificar el próximo salto inmediato de dirección IPv6. Este próximo salto está especificado por un RTE especial, en la ruta del próximo salto de la tabla de entrada. El siguiente salto (Next Hop) RTE está identificado por el valor de 0xFF en el campo métrica. El campo prefijo (Prefix) especifica la dirección IPv6 del próximo salto; etiqueta de ruta y longitud de prefijo son configurados a cero en la transmisión e ignorados en la recepción.

El próximo salto de la Tabla de Entrada de Ruteo (RTE) tiene el siguiente formato:

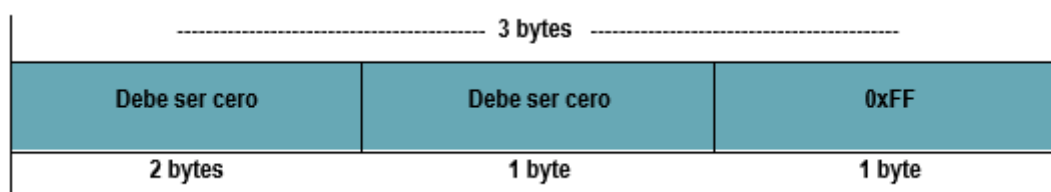


Figura 80: Próximo salto RTE

Especificando un valor de 0:0:0:0:0:0:0:0 en el campo prefijo de un próximo salto RTE indica que la próxima dirección del salto debe ser el originador del anuncio RIP. Una dirección especificada como un próximo salto debe tener una dirección local de enlace.

El propósito del próximo salto RTE es eliminar los paquetes que están siendo ruteados a través de saltos extras en el sistema. Esto es particularmente útil cuando RIPng no está corriendo sobre todos los routers en una red. Es de hacer notar que el próximo salto RTE es "advisory". Esto es, si la información provista es ignorada, una posibilidad sub-óptima, pero absolutamente válida, la ruta puede ser tomada. Si la dirección del próximo salto recibida no es una dirección de enlace local, esta debe ser tratada como 0:0:0:0:0:0:0:0.

RIPng no es ni más ni menos potente que RIP, pero proporciona una manera sencilla de crear una red IPv6 sin necesidad de crear un nuevo protocolo de enrutamiento.

7.3 Características de RIPng:

- Basado en IPv4 RIP versión 2 (RIPv2) y es similar a RIPv2.
- Usa IPv6 para el transporte.
- Incluye el prefijo IPv6 y la dirección IPv6 del siguiente salto.
- Usa el grupo Multicast FF02::9 como dirección de destino para las actualizaciones de RIP (similar a la función de Broadcast que realiza RIP en IPv4).
- Envía actualizaciones por el puerto UDP 521



La versión 1 de RIPng soporta 2 comandos: Request y Response. Un Request es utilizado para preguntar por toda o parte de la tabla de ruteo. En muchos casos, los Request son enviados como multicasts desde el puerto RIPng (521). Si la información para solo un router es necesitada, esa solicitud sería enviada directamente a ese router desde un puerto que no sea el puerto RIPng. Hay 3 tipos de Response: una respuesta a una consulta específica; una actualización regular, que es una respuesta no solicitada enviada cada 30 segundos a todos los routers vecinos; y una actualización accionada causada por un cambio de ruta.

RIPng vs RIPv2

RIPng establece que el protocolo utiliza iguales conceptos y convenciones de las especificaciones originales de RIPv1, además de algunos conceptos de RIPv2.

Características	RIPv2	RIPng
Anuncio de Rutas	IPv4	IPv6
Mensajes RIP usan protocolos de la capa 3 y 4	IPv4:UDP	IPv6:UDP
Puerto UDP	520	521
Usa Vector Distancia	Si	Si
Distancia Administrativa por Defecto	120	120
Soporta VLSM	Si	Si
Sumarización Automática	Si	No
Split Horizont (horizonte Dividido)	Si	Si
Envenenamiento en Reversa	Si	Si
Actualización periódicas completas cada 30 Seg	Si	Si
Actualizaciones Disparadas	Si	Si
Conteo de Saltos	Si	Si



Métrica infinita	16	16
Soporta Rutas etiquetadas	Si	Si
Direcciones multicast de Actualizaciones	224.0.0.9	FF02::9
Usa Autenticación	RIP specific	IPv6AH/ESP

Tabla 13: RIPng vs RIPv2

Algunas diferencias se refieren específicamente a IPv6.

- Primero, los mensajes propios lista los prefijos y longitudes IPv6, En lugar de máscaras y subredes.
- En RIPv1 y 2, RIP encapsula mensajes de actualización de RIP dentro de un encabezado en IPv4 con UDP; con IPv6, la encapsulación se utiliza para paquetes IPv6, de nuevo con un encabezado UDP.
- Algunas pequeñas diferencias en el formato de mensaje de actualización también existen, con la diferencia obvia que en la lista de actualizaciones IPv6 van los prefijos y la longitud del prefijo.
- Otra diferencia es que IPv6 soporta la autenticación usando IPsec como encabezado de autenticación (AH), RIPng no admite de forma nativa la autenticación compatible, sino que depende de IPsec.

7.4 Configuración RIPng

RIPng usa un estilo nuevo de comandos para la configuración básica, pero la mayoría tiene características opcionales y los comandos de verificación se parecen mucho a los utilizados en RIP para IPv4.

La gran diferencia entre RIPv2 y la configuración RIPng es que RIPng descarta el comando Network de RIP a diferencia del subcomando de la interfaz **ipv6 rip nombre enable**, que permite que RIPng se active para la interfaz. Otra diferencia se relaciona con el enrutamiento de IPv4 e IPv6: las rutas de IPv4 con un IOS por defecto (tiene el comando ip routing), pero estos IOS no establecen una ruta IPv6 por defecto (no ipv6 unicast routing).



Ejemplo de una implementación de RIPng (RIP New Generation) para el enrutamiento.

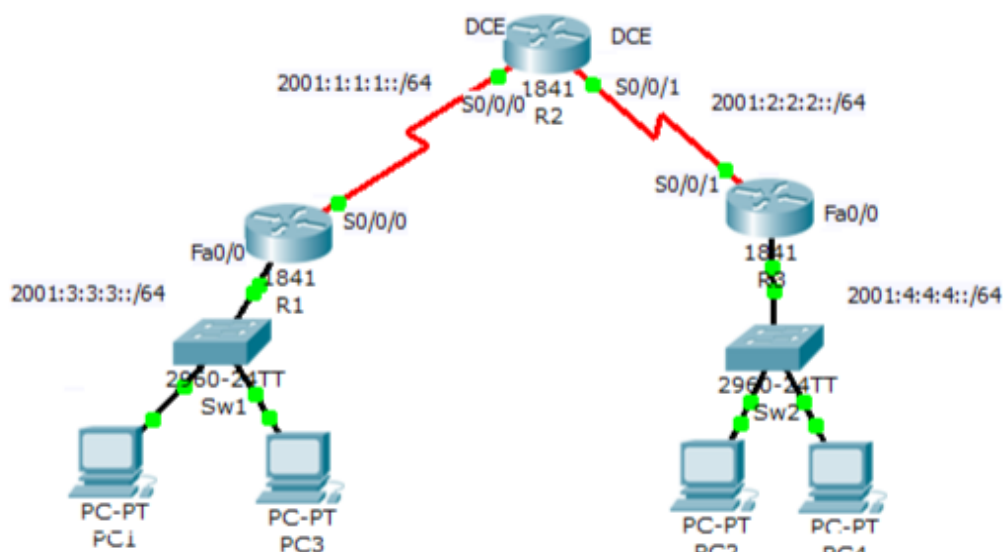


Figura 81: Topología de red con RIPng

Objetivos

- Conectar una red de acuerdo con el Diagrama de topología.
- Eliminar la configuración de inicio y recargar un router para volver al estado predeterminado.
- Habilitar IPv6 unicast routing.
- Configurar y activar direcciones IPv6.
- Autorizar RIPng en las interfaces correspondientes.
- Probar y verificar las configuraciones.
- Reflexionar sobre la implementación de la red y documentar el procedimiento



Tabla de Direcccionamiento

Dispositivo	Interfaz	Dirección IP	Máscara de subred	Gateway predeterminado
R1	Fa0/0	2001:3:3:3::1	/64	N/A
	S0/0/0	2001:1:1:1::1	/64	N/A
R2	S0/0/0	2001:1:1:1::2	/64	N/A
	S0/0/1	2001:2:2:2::1	/64	N/A
R3	Fa0/0	2001:4:4:4::1	/64	N/A
	S0/0/1	2001:2:2:2::2	/64	N/A
PC1	NIC	2001:3:3:3::2	/64	2001:3:3:3::1
PC2	NIC	2001:4:4:4::2	/64	2001:4:4:4::1
PC3	NIC	2001:3:3:3::3	/64	2001:3:3:3::1
PC4	NIC	2001:4:4:4::3	/64	2001:4:4:4::1

Tabla 14: Ejemplo, tabla de direccionamiento con RIPng

Configurar direcciones IPv6 en los routers

```
Router(config)# ipv6 unicast-routing
```

Configurar la dirección IPv6 y RIPng en las Interfaces de R1

```
R1(config)# interface fa0/0
R1(config-if)# ipv6 address 2001:3:3:3::1/64
R1(config-if)# ipv6 rip RUTEO enable
R1(config-if)# no shutdown
R1(config-if)# end
```

En la interfaz S0/0/0 del router 1 (R1), configure la dirección IPv6 que se encuentra en la tabla de direccionamiento, habilite el proceso de RIPng llamado “RUTEO”. Configure la dirección IP que se encuentra en la tabla de direccionamiento de la interfaz FastEthernet0/0 en R1.

```
R1(config)# interface s0/0/0
R1(config-if)# ipv6 address 2001:1:1:1::1/64
R1(config-if)# clock rate 64000
R1(config-if)# ipv6 rip RUTEO enable
R1(config-if)# no shutdown
R1(config-if)# end
```

Configurar la dirección IPv6 y RIPng en R2

En la interfaz S0/0/0 del router 2 (R2), configure la dirección IPv6 que se encuentra en la tabla de direccionamiento, habilite el proceso de RIPng llamado “RUTEO”

```
R2(config)# interface s0/0/0
R2(config-if)# ipv6 address 2001:1:1:1::2/64
R2(config-if)# clock rate 64000
R2(config-if)# ipv6 rip RUTEO enable
R2(config-if)# no shutdown
R2(config-if)# end
```



Configure la dirección IP que se encuentra en la tabla de direccionamiento de la interfaz Serial0/0/1 en R2

```
R2(config)# interface s0/0/1
R2(config-if)# ipv6 address 2001:2:2::1/64
R2(config-if)# clock rate 64000
R2(config-if)# ipv6 rip RUTEO enable
R2(config-if)# no shutdown
R2(config-if)# end
```

Configurar la dirección IPv6 y RIPng en R3

En la interfaz S0/0/1 del router 3 (R3), configure la dirección IPv6 que se encuentra en la tabla de direccionamiento, habilite el proceso de RIPng llamado “RUTEO” y configure la señal de reloj de 64000 utilizando los siguientes comandos:

Configure la dirección IP que se encuentra en la tabla de direccionamiento de la interfaz FastEthernet0/0 en R3.

```
R3(config)# interface fa0/0
R3(config-if)# ipv6 address 2001:4:4:4::1/64
R3(config-if)# ipv6 rip RUTEO enable
R3(config-if)# no shutdown
R3(config-if)# end
```

NOTA: Para habilitar RIPng solamente se debe ingresar a la interfaz de router que se desea publicar en el proceso RIP ingresar el comando `ipv6 rip IDENTIFICADOR enable` donde “IDENTIFICADOR” es un ID de proceso. El nombre de proceso no debe ser el mismo en todos los routers. Como es un identificador local, puede ser el mismo o diferente en toda la red RIP y el enrutamiento funcionará igualmente. Sin embargo, es necesario que en el router, todas las interfaces pertenezcan al mismo proceso.

Si deseamos verificar que todas las rutas han sido publicadas satisfactoriamente se debe de usar el comando **show ipv6 route**.



8. EIGRPv6

El propietario cisco desarrollo **Enhanced Interior Gateway Protocol** (EIGRP) para reducir la brecha entre los protocolos de vector distancia tradicionales (IGRP, RIP) y los protocolos de enlace-estado avanzados (OSPF, IS-IS). Integra algunas de las capacidades de este último para mejorar el funcionamiento y la escalabilidad de los anteriores.

Sin embargo, la intención era evitar algunas de las restricciones topológicas que a veces se asocian a los protocolos de enlace-estado. El resultado es un simple rápido-convergente, flexible y escalable protocolo de enrutamiento, que en gran medida se adopta en muchas redes de la empresa y en el borde de algunas redes de ISP.

EIGRPv6 es una modificación al protocolo original EIGRP .por el cual se esperaba mucho, así que fue introducido en Cisco IOS versión 12.4 (6) T. Ofrece las siguientes mejoras:

- Algoritmo de actualización por difusión (DUAL) puede determinar si una ruta anunciada por un vecino es libre de bucles y para identificar las rutas alternativas sin esperar a las actualizaciones de otros routers.
- Éste consulta activamente a los vecinos cuando los destinos se vuelven inalcanzables, y que conduce para tiempos de convergencia competitivos.
- Utiliza los paquetes de saludo (hello) para mantener el estado de su vecino, lo que lleva a una convergencia más rápida.
- Utiliza protocolo de transporte fiable para el intercambio de las actualizaciones y si elimina la necesidad de actualizaciones periódicas completas.
- Utiliza métricas complejas que proporcionan flexibilidad en la selección de la ruta.

EIGRPv6 utiliza un nuevo módulo de Protocolo Dependiente (PDM) y transporta IPv6, tiene tres nuevos valores Tipo Longitud (TLV):

- IPv6_REQUEST_TYPE (0x0401)
- IPv6_METRIC_TYPE (0x0402)
- IPv6_EXTERIOR_TYPE (0x0403)

Estos elementos de información EIGRP contienen cada uno un tipo que identifica el tipo de Información, la longitud de la información en bytes, y el valor que el parámetro contiene; de ahí el nombre Tipo Longitud Valor (TLV).

Usa el código Hello Message proveniente de la dirección de enlace local y destinados a la dirección IPv6 link-local, alcance de dirección Multicast FF02::A (todos los routers EIGRP).Aquí "no autosummary" esta



desactivado y no se necesita para IPv6 porque cada ruta utiliza máscara de longitud variable (VLSM) por defecto. EIGRPv6 no usa horizonte dividido (Split horizon) porque IPv6 soporta múltiples prefijos por interfaz.

Históricamente EIGRP ha tenido problemas de seguridad. Hay conjuntos de herramientas disponibles que podrían atacar routers EIGRP y el hecho de que se utilicen comunicaciones Multicast entre vecinos e intercambio de información en texto claro. Los ataques también aparecen en el uso de mensajes de despedida EIGRP. Sin embargo, estas vulnerabilidades fueron mitigadas mediante el uso estáticamente configurado de vecinos adyacentes y mediante la autenticación MD5. Usando anuncios estáticos de vecinos, EIGRP entonces fuerza mensajes unicast Hello Message en lugar de multidifusión. Sin embargo, todavía podría haber problemas con Spoofing unicast mensajes de saludo, a pesar de que se han usado estas técnicas.

8.1 Diferencias entre EIGRP IPv6 e IPv4

EIGRP para IPv4 e IPv6 tiene fuertes similitudes. Existen unas pocas diferencias debido a algunos aspectos específicos de IPv6:

- El ID del router en el proceso EIGRP mantiene longitud de 32 bits. Se deriva de una dirección IPv4 que se encuentra en una de las interfaces configuradas o se configura manualmente.
- En un sólo router con IPv6, el proceso EIGRP no se inicia hasta que el ID está configurado manualmente. La dirección de origen (SA) del EIGRP Hello es el vínculo de la dirección local de la interfaz de transmisión; la dirección de destino (DA) es FF02 :: A (todos los routers EIGRP, linkscope dirección de multidifusión).
- El formato del paquete Hello implica que dos routers vecinos no tienen que compartir el mismo prefijo en el enlace para ver Hello del otro. Los paquetes enviados a los pares específicos por unidifusión, en cuyo caso compartiendo el mismo prefijo en el enlace se vuelven relevantes.
- La autenticación de EIGRP para IPv6 se basa en la incorporación de características de autenticación y confidencialidad propias de IPv6, en vez del algoritmo MD5 para la autenticación que usa en IPv4. EIGRP para IPv6 encapsula sus mensajes en paquetes IPv6, en lugar de paquetes IPv4.
- EIGRP para IPv4 por defecto utiliza automáticamente sumarización de rutas en los límites de las redes IPv4 con clase; IPv6 no tiene ningún concepto de las redes con clase, por lo que EIGRP para IPv6 no puede realizar ninguna sumarización automática.
- EIGRP para IPv6 anuncia prefijos y longitudes para IPv6, en lugar de información de las máscaras de subred de IPv4.
- A diferencia de EIGRP para IPv4, no hay horizonte dividido en EIGRP para IPv6, porque en IPv6 múltiples prefijos podrían estar presentes en la misma interfaz de un router.
- EIGRP en IPv6 no requiere vecinos en la misma subred IPv6 como requisito para convertirse en vecinos.



- EIGRP para IPv6 estará habilitada para operar dentro de redes privadas virtuales (VPN) de una manera similar a como EIGRP para IPv4.
- Al igual que con EIGRP para IPv4, EIGRPv6 utiliza el número de protocolo 88. Sin embargo para EIGRPv6, usa una cabecera de extensión IPv6 con el valor 88.

Características	EIGRP-IPv4	EIGRP-IPv6
Anuncia Rutas	IPv4	IPv6
Protocolo de capa 3 EIGRP con mensaje	IPv4	IPv6
Tipo de encabezado del protocolo capa 3	88	88
Puertos UDP	No Aplica	No Aplica
Uso de sucesor, y sucesor factible	Si	Si
Uso de DUAL	Si	Si
Soporta VLSM	Si	Si
Sumarización Automática	Si	No Aplica
Actualizaciones disparadas	Si	Si
Métrica compuesta, por defecto ancho de banda y retardo	Si	Si
Métrica significado Infinito	$2^{32}-1$	$2^{32}-1$
Soporta etiquetado de rutas	Si	Si
Dirección Actualización Multicast	224.0.0.10	FF02::10
Autenticación	EIGRP Specific	IPv6 AH/ESP

Tabla 15: Características de .EIGRPv4 vs EIGRPv6.

8.2 Ejemplo de configuración de EIGRPv6.

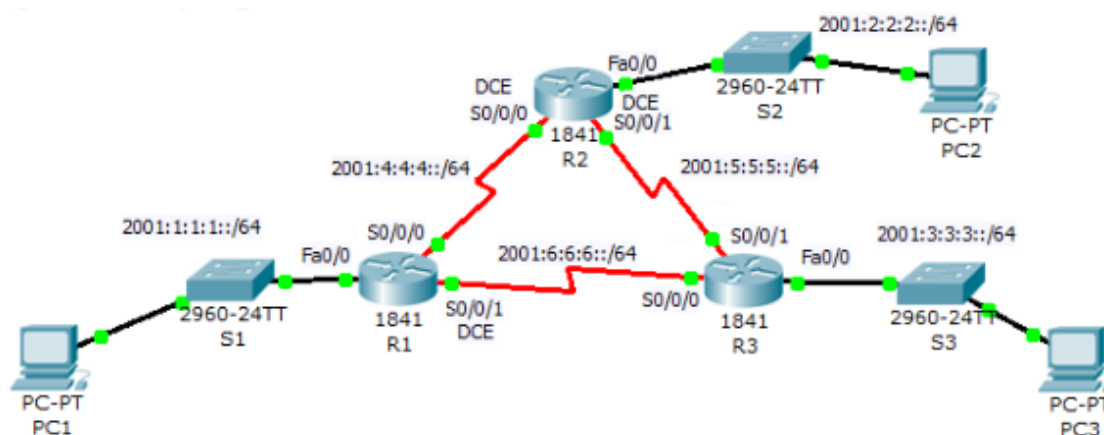


Figura 82: Topología de EIGRP



Habilitar IPv6 unicast routing

Todos los router deben tener habilitado el soporte IPv6 unicast routing, para habilitar RIPng.

Tabla de Direccionamiento

Dispositivo	Interfaz	Dirección IP	Máscara de subred	Gateway predeterminado
R1	Fa0/0	2001:1:1:1::1	/64	N/A
	S0/0/0	2001:4:4:4::1	/64	N/A
	S0/0/1	2001:6:6:6::1	/64	N/A
R2	Fa0/0	2001:2:2:2::1	/64	N/A
	S0/0/0	2001:4:4:4::2	/64	N/A
	S0/0/1	2001:5:5:5::1	/64	N/A
R3	Fa0/0	2001:3:3:3::1	/64	N/A
	S0/0/0	2001:6:6:6::2	/64	N/A
	S0/0/1	2001:5:5:5::2	/64	N/A
PC1	NIC	2001:1:1:1::2	/64	2001:1:1:1::1
PC2	NIC	2001:2:2:2::2	/64	2001:2:2:2::1
PC3	NIC	2001:3:3:3::2	/64	2001:3:3:3::1

Tabla 16: Ejemplo, tabla de direccionamiento con EIGRPv6

Configurar la dirección IPv6 y EIGRPv3 en R1

En la interfaz S0/0/0 del router 1 (R1), configure la dirección IPv6 que se encuentra en la tabla de direccionamiento, habilite el enrutamiento dinámico de EIGRPv3 en la interfaz, utilizando los siguientes comandos:

```
R1(config)# interface s0/0/0
R1(config-if)# ipv6 address 2001:4:4:4::1/64
R1(config-if)# ipv6 elgrp 1
R1(config-if)# no shutdown
R1(config-if)# end
```

En la interfaz S0/0/1 del router 1 (R1), configure la dirección IPv6 que se encuentra en la tabla de direccionamiento, habilite el enrutamiento dinámico de EIGRPv3 en la interfaz y configure la señal de reloj de 64000 utilizando los siguientes comandos:

```
R1(config)# interface s0/0/1
R1(config-if)# ipv6 address 2001:6:6:6::1/64
R1(config-if)# clock rate 64000
R1(config-if)# ipv6 elgrp 1
R1(config-if)# no shutdown
R1(config-if)# end
```



Configure la dirección IP que se encuentra en la tabla de direccionamiento de la interfaz FastEthernet0/0 en (R1).

```
R1(config)# interface fa0/0
R1(config-if)# ipv6 address 2001:1:1:1::1/64
R1(config-if)# ipv6 eigrp 1
R1(config-if)# no shutdown
R1(config-if)# end
```

Para configurar EIGRPv3 en el router se utilizan los siguientes comandos:

```
R1(config)# ipv6 router eigrp 1
R1(config-rtr)# router-id 10.1.1.1
R1(config-rtr)# no shutdown
```

Configurar la dirección IPv6 y EIGRPv3 en R2

En la interfaz S0/0/0 del router 2 (R2), configure la dirección IPv6 que se encuentra en la tabla de direccionamiento, habilite el enrutamiento dinámico de EIGRPv3 en la interfaz, utilizando los siguientes comandos:

```
R2(config)# interface s0/0/0
R2(config-if)# ipv6 address 2001:4:4:4::2/64
R2(config-if)# clock rate 64000
R2(config-if)# ipv6 eigrp 1
R2(config-if)# no shutdown
R2(config-if)# end
```

Configure la dirección IP que se encuentra en la tabla de direccionamiento de la interfaz Serial0/0/1 en (R2)

```
R2(config)# interface s0/0/1
R2(config-if)# ipv6 address 2001:5:5:5::1/64
R2(config-if)# clock rate 64000
R2(config-if)# ipv6 eigrp 1
R2(config-if)# no shutdown
R2(config-if)# end
```

Para configurar EIGRPv3 en el router se utilizan los siguientes comandos

```
R2(config)# ipv6 router eigrp 1
R2(config-rtr)# router-id 10.1.1.1
R2(config-rtr)# no shutdown
```

Configurar la dirección IPv6 y EIGRPv3 en R3



En la interfaz S0/0/0 del router 3 (R3), configure la dirección IPv6 que se encuentra en la tabla de direccionamiento, utilizando los siguientes comandos:

```
R3(config)# interface s0/0/0  
R3(config-if)# ipv6 address 2001:6:6:6::2/64  
R3(config-if)# ipv6 eigrp 1  
R3(config-if)# no shutdown  
R3(config-if)# end
```

Configure la dirección IP que se encuentra en la tabla de direccionamiento de la interfaz Serial0/0/1 en (R3)

```
R3(config)# interface s0/0/1  
R3(config-if)# ipv6 address 2001:5:5:5::2/64  
R3(config-if)# ipv6 eigrp 1  
R3(config-if)# no shutdown  
R3(config-if)# end
```

Configure la dirección IP que se encuentra en la tabla de direccionamiento de la interfaz FastEthernet0/0 en (R3).

```
R3(config)# interface fa0/0  
R3(config-if)# ipv6 address 2001:3:3:3::1/64  
R3(config-if)# ipv6 eigrp 1  
R3(config-if)# no shutdown  
R3(config-if)# end
```

Para configurar EIGRPv3 en el router se utilizan los siguientes comandos:

```
R3(config)# ipv6 router eigrp 1  
R3(config-rtr)# router-id 10.1.1.1  
R3(config-rtr)# no shutdown
```

NOTA: Además de indicar la interfaz donde usará el protocolo EIGRP, se requiere del identificador para anunciar las redes que están conectadas directamente a los routers.



9. OSPFv3

OSPF son las siglas **Open Shortest Path First** (El camino más corto primero), un protocolo de enrutamiento jerárquico de pasarela interior o IGP (Interior Gateway Protocol), que usa el algoritmo SmoothWall Dijkstra enlace-estado (LSE-Link State Algorithm) para calcular la ruta más idónea.

Su medida de métrica se denomina Cost, y tiene en cuenta diversos parámetros tales como el ancho de banda y la congestión de los enlaces. OSPF construye además una base de datos enlace-estado (link-state database, LSDB) idéntica en todos los enrutadores de la zona.

OSPF puede operar con seguridad usando MD5 para autenticar sus puntos antes de realizar nuevas rutas y antes de aceptar avisos de enlace-estado.

OSPFV3 es el protocolo de encaminamiento OSPF para IPv6, además fue publicado en el RFC 2740, fue basada en OspfV2 y modificada para que soportara distribuir prefijos IPv6. Usa IPv6 como transporte aunque tiene el mismo nombre que OspfV2, son protocolos diferentes.

La implementación del protocolo para IPv6 incluye estas características:

➤ **Similar a IPv4.**

Los mismos mecanismos; pero mejor reescritura en los protocolos internos.

➤ **Nuevas características de IPv6:**

- Toda la semántica de IPv4 específica es removida.
- Los ISP con direcciones de IPv6.
- Las direcciones link-local se usan en el origen.
- Transporta IPv6.
- OSPF para IPv6 es un estándar específico de la IETF.

➤ **Basada en la versión 2 (OSPFv2), con mejoras:**

- Distribuye prefijos IPv6.
- Se ejecuta directamente en IPv6.

➤ **Esta aplicación agrega estos atributos específicos a IPv6:**

- Direcciones de 128 bits.
- Dirección de enlace local.
- Múltiples direcciones y peticiones por interfaces.
- Autenticación (ahora usa IPsec).



- OSPFv3 corre en un enlace en lugar de una subred.

El estado de enlace es una descripción de la interfaz y su relación con sus dispositivos de red vecinos. La información de interfaz incluye el prefijo IPv6 de la interfaz, la máscara de red, el tipo de red a la que se conecta, los routers conectados a la red, y así sucesivamente.

Esta información se propaga en diversos tipos de anuncios de estado de enlace (link-state Advertisement - LSAs). Una colección de datos LSA en un Router es almacenado en una base de datos de estado de enlace (link-state database - LSDB). El contenido de las bases de datos, cuando es sometido al algoritmo Dijkstra, resulta en la creación de la tabla de enrutamiento OSPF.

La diferencia entre la base de datos y la tabla de enrutamiento es que la base de datos contiene una completa colección de datos puros. La tabla de enrutamiento contiene una lista más corta de caminos para conocer los destinos específicos a través de puertos de interfaz del Router. OSPFv3, que se describe en RFC 2740, soporta IPv6.

La operación de OSPFv3 conforme el RFC soportan los modos multiacceso de **no difusión** (NBMA) y topología punto a multipunto, también soporta otros modos de Cisco, como punto a punto y Broadcast, incluyendo la interfaz.

OSPFv3 utiliza los mismos tipos de paquetes básicos como OSPFv2, como los paquetes hello, descripción de la base de datos (también llamada paquete de descripción de la base de datos), petición de estado de enlace (LSR), actualización de estado de enlace (LSU), y LSA.

9.1 Tipos de paquetes OSPF

OSPFv3 tiene los mismos 5 tipos de paquetes, pero algunos campos son cambiados.

Tipo de paquete	Descripción
1	HELLO
2	Descripción de la base de datos.
3	Requerimiento del estado del enlace.
4	Actualización del estado del enlace.
5	Acuse de recibo del estado del enlace.

Tabla 17: Tipos de paquetes OSPF.

Todos los paquetes OSPFv3 tienen 16 bytes de encabezado versus 24 bytes de encabezado de OSPFv2.



Version	Tipo	Long. Paquete	Version	Tipo	Long. Paquete
ID Router			ID Router		
ID Area			ID Area		
Checksum		T. Autenticación	Checksum		ID Instancia 0
Autenticación					
Autenticación					

Tabla 18: Encabezado OSPFv3 vs OSPFv2

9.2 Diferencias entre OSPFv2 y OSPFv3

Las diferencias entre OSPFv2 y OSPFv3 incluyen lo siguiente:

OSPFv3 es un protocolo de procesamiento por enlace, no por subred.

- IPv6 conecta las interfaces de los enlaces.
- Múltiples subredes IPv6 se pueden asignar a un simple enlace.
- Dos nodos se pueden conectar a un simple enlace directamente y no compartir una subred en común.
- El término "red" y "subred" están comenzando a reemplazarse por "enlace".
- Las interfaces en OSPF ahora se conectan a los enlaces instalados en la subred.

OSPF para IPv6 se ejecuta en un enlace en lugar del comportamiento de IPv4 por subred IP. IPv6 usa el término "vínculo" para indicar "una facilidad de comunicación o medio sobre el cual los nodos pueden comunicarse en la capa de enlace." Por lo tanto, los términos "red" y "subred", utilizada en la especificación OSPF IPv4 se sustituyen por "vínculo".

La declaración Network en el modo de subcomando del router OSPFv2 es reemplazado por el comando de interfaz `ipv6 ospf process-id area area-id [instance instance-id]`.

Se usan direcciones de Enlace Local

OSPFv3 utilizan direcciones de enlace local IPv6 para identificar las adyacencias de vecinos OSPFv3. Por lo tanto, cuando se configura el comando **IPv6 ospf Neighbor**, la dirección **IPv6** que se debe utilizar es la dirección local de vínculo del vecino.

9.3 Soporte para múltiples instancia OSPFv3

Separa los sistemas autónomos, cada uno de OSPF en ejecución, utilizan un vínculo común. Un único enlace puede pertenecer a múltiples áreas.



OSPFv3 utiliza un nuevo campo, llamado ID (identificador) de la instancia, para permitir múltiples instancias por enlace. Para tener 2 instancias de conversaciones unos con otros, deben compartir la misma instancia ID. Por defecto, el identificador de instancia se establece en 0.

Direcciones Multicast

- Direcciones multicast:
 - FF02::5 representa todos los routers SPF en un enlace local; equivalente a la 224.0.0.5 de OSPFv2
 - FF02::6 representa todos los routers DR en un enlace local; equivalente a la 224.0.0.6 de OSPFv2
- Semántica de la dirección removida:
 - La dirección IPv6 no es más grande que el encabezado del paquete ospf (la parte de la información de carga útil)
 - Los routers LSA y redes LSA no llevan direcciones IPv6
 - El router ID, el área ID, y el estado del enlace ID son de 32 bits
 - El DE y el BDR están identificados con el router ID y por la longitud de la dirección IP
- Seguridad:
 - OSPFv3 usa para IPv6 AH y ESP como extensiones de los encabezados instalados en varios mecanismos definidos en OSPFv2

Tabla 19: Direcciones Multicast con OSPF

- **FF02::5** representa primero la ruta más corta de todos los routers (shortest Path First - SPF) en el alcance de enlace local, equivalente a 224.0.0.5 en OSPFv2.
- **FF02::6** representa todos los routers designados (designated routers - DRs) en el alcance de enlace local, equivalente a 224.0.0.6 en OSPFv2.

Supresión de la semántica de dirección

- Las direcciones IPv6 no están más presentes en el encabezado del paquete OSPF (parte de la información de carga útil).
- Los routers LSA y las redes LSA no llevan direcciones IPv6.
- El Router ID, área ID, y el ID de enlace permanecen en 32 bits.
- El Router DR y el router de respaldo designado (backup designated router - BDR) son identificados por su router ID y no por su dirección IP.

9.4 Seguridad

OSPFv3 utiliza el encabezado de autenticación IPv6 (Authentication Header - AH) y los encabezados de extensión de seguridad de carga útil (Encapsulating Security Payload - ESP), en lugar de la variedad de mecanismos definidos en OSPFv2.



La autenticación ya no forma parte de OSPF. Ahora es el trabajo de IPv6 asegurar que el nivel correcto de autenticación este en uso.

9.4 Tipos de LSA para IPv6

LSA OSPFv3 incluyen las siguientes Las características:

- El LSA se compone de un Router ID, área ID, e ID de estado de enlace. Son Cada uno de 32 bits. Aunque, estén escritos en notación decimal, no son derivados de una dirección IPv4.
- Las LSAs de router y LSAs de red contienen solo Id de 32 bits. No contienen direcciones.
- Las LSAs tienen alcance de transmisión que define el diámetro hacia donde deben transmitir:
 - **Enlace local:** Fluye a todos los routers en la red.
 - **Área:** Fluye a todos los routers dentro de un área OSPF.
 - **Sistema Autónomo:** Fluye a todos los routers en un sistema autónomo entero OSPF.

OSPFv3 soporta el envío de LSAs desconocidos basado en el alcance de transmisión. Esto puede ser útil en una NSSA.

OSPFv3 toma las ventajas del Multicasting IPv6, usando FF02::5 para todos los routers OSPF, y FF02::6 para el DR y BDR OSPF.

Los dos LSAs renombrados son los siguientes:

LSAs de prefijo Interárea para los routers de área de frontera (ABRs) (tipo 3): Las LSA de tipo 3 anuncian las redes internas de los routers en otras áreas (rutas interárea). Las LSAs tipo 3 pueden representar una única red o un grupo de redes sumariadas en una advertencia. Solo los ABRs generan LSAs sumariadas. En OSPF para IPv6, las direcciones para estas LSAs son expresadas como prefijos, prefijos de longitud en vez de dirección, máscara. La ruta por defecto se expresa como un prefijo con longitud 0.

LSAs de Router Interárea para Sistemas Autónomos de Routers de Frontera (ASBRs) (tipo 4): Las LSAs tipo 4 advierten la ubicación de un ASBR. Los routers que tratan de alcanzar una red externa usan estas advertencias para determinar la mejor ruta al siguiente salto. Los ASBRs generan LSAs tipo 4.

Las dos nuevas LSAs en IPv6 son las siguientes:

LSAs de enlace (tipo 8): Las LSAs de tipo 8 tienen un alcance de enlace local y nunca se transmiten más allá del enlace con el cual están asociados. Las LSAs de enlace proveen la dirección de enlace local del router para todos los demás routers ligados al enlace. Las LSAs de enlace también informan a los otros routers ligados al enlace de una lista de prefijos IPv6 para asociarlos con el enlace, y permitir al router afirmar una colección de bits opcionales para asociarlo con la red LSA que se originada por el enlace.



LSAs de prefijo Intra-Área (tipo 9): Un router puede originar múltiples LSAs de prefijo intra-área para cada Router o red de tránsito, cada uno con un único ID de estado de enlace. El ID de estado de enlace para cada prefijo intra-área LSA describe su asociación a la LSA de Router o LSA de red. El ID de estado de enlace también contiene prefijos para redes stub y de tránsito.

	LSA Function Code	LSA type
Router-LSA	1	0x2001
Network-LSA	2	0x2002
Inter-Area-Prefix-LSA	3	0x2003
Inter-Area-Router-LSA	4	0x2004
AS-External-LSA	5	0x2005
Group-membership-LSA	6	0x2006
Type-7-LSA	7	0x2007
Link-LSA	8	0x2008
Intra-Area-Prefix-LSA	9	0x2009

Tabla 20: LSAs (tipo9)

Prefijo de Dirección y LSAs

- Un prefijo de dirección se produce en casi todas las nuevas LSAs definidas.
- El prefijo es representado por tres campos: prefijo de longitud, prefijo de opciones, y prefijo de dirección. En OSPF para IPv6, las direcciones para las LSAs son expresadas como prefijos, prefijo de longitud en vez de dirección, máscara.
- La ruta por defecto se expresa como un prefijo con longitud 0.
- Las LSAs de tipo 3 y 9 llevan toda la información del prefijo IPv6, la cual, en IPv4 se incluye en las LSAs de router y de red.

9.5 Configuración de OSPFv3 en IPv6.

Muchos comandos OSPFv3 son similares a los de OSPFv2. En muchos casos, simplemente debe reemplazar cada prefijo **ip** en el comando OSPF con **ipv6**. Por ejemplo, en vez de usar el comando **ip Address** para asignar una dirección IPv6 use el comando **ipv6 Address**. Para ver las rutas IPv6, use el comando de emisión **show ipv6 route**.

La configuración de OSPFv3 no es un modo de subcomando del comando **router ospf** como en la configuración de OSPFv2. Por ejemplo, en vez de utilizar el comando **Network area** para identificar redes que forman parte de la red OSPFv3, las interfaces se configuran directamente para especificar que las redes IPv6 son parte de la red OSPFv3.



Ejemplo de una implementación con OSPFv3

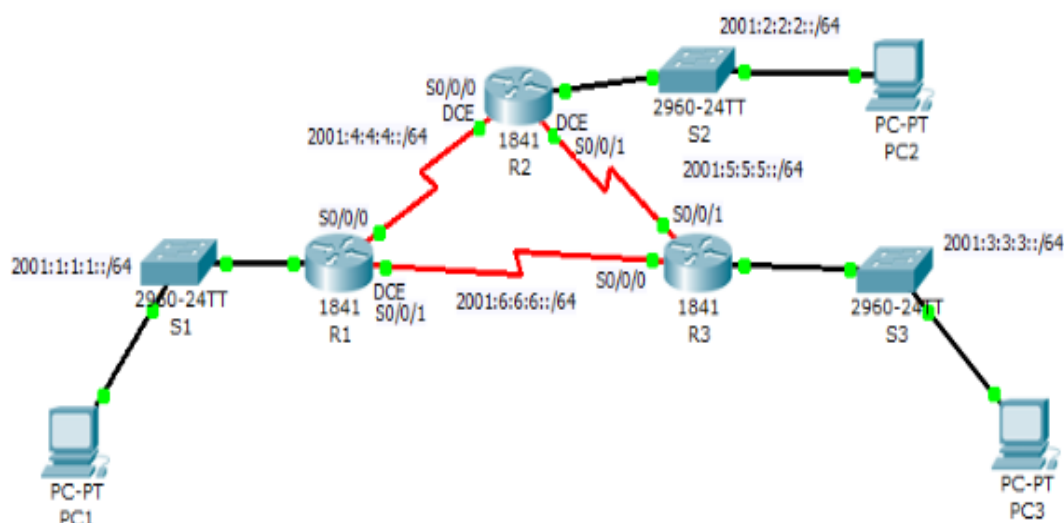


Figura 83: Topología de OSPFv3

Tabla de Direccionamiento

Dispositivo	Interfaz	Dirección IP	Máscara de subred	Gateway predeterminado
R1	Fa0/0	2001:1:1:1::1	/64	N/A
	S0/0/0	2001:4:4:4::1	/64	N/A
	S0/0/1	2001:6:6:6::1	/64	N/A
R2	Fa0/0	2001:2:2:2::1	/64	N/A
	S0/0/0	2001:4:4:4::2	/64	N/A
	S0/0/1	2001:5:5:5::1	/64	N/A
R3	Fa0/0	2001:3:3:3::1	/64	N/A
	S0/0/0	2001:6:6:6::2	/64	N/A
	S0/0/1	2001:5:5:5::2	/64	N/A
PC1	NIC	2001:1:1:1::2	/64	2001:1:1:1::1
PC2	NIC	2001:2:2:2::2	/64	2001:2:2:2::1
PC3	NIC	2001:3:3:3::2	/64	2001:3:3:3::1

Tabla 21: Tabla de direcciones con OSPFv6.

Habilitar IPv6 unicast routing

Todos los router deben tener habilitado el soporte IPv6 unicast routing, para habilitar RIPng.

```
Router(config)# ipv6 unicast-routing
```

Configurar la dirección IPv6 y OSPF en R1

En la interfaz s0/0/0 de R1, configure la dirección IPv6 que se encuentra en la tabla de dirección IPv6 que se encuentra en la tabla de direccionamiento, habilite el enrutamiento dinámico de OSPF en la interfaz, utilizando los siguientes comandos:



```
R1(config)# interface s0/0/0
R1(config-if)# ipv6 address 2001:4:4:4::1/64
R1(config-if)# ipv6 ospf 1 area 0
R1(config-if)# no shutdown
R1(config-if)# end
```

En la interfaz S0/0/1 del router 1 (R1), configure la dirección IPv6 que se encuentra en la tabla de direccionamiento, habilite el enrutamiento dinámico de OSPFv3 en la interfaz y configure la señal de reloj de 64000 utilizando los siguientes comandos:

```
R1(config)# interface s0/0/1
R1(config-if)# ipv6 address 2001:6:6:6::1/64
R1(config-if)# clock rate 64000
R1(config-if)# ipv6 ospf 1 area 0
R1(config-if)# no shutdown
R1(config-if)# end
```

Configure la dirección IP que se encuentra en la tabla de direccionamiento de la interfaz FastEthernet0/0 en R1.

```
R1(config)# interface fa0/0
R1(config-if)# ipv6 address 2001:1:1:1::1/64
R1(config-if)# ipv6 ospf 1 area 0
R1(config-if)# no shutdown
R1(config-if)# end
```

Configurar OSPFv3 en el router se utilizan los siguientes comandos:

```
R1(config)# ipv6 router ospf 1
R1(config-rtr)# router-id 10.1.1.1
```

Configurar la dirección IPv6 y OSPFv3 en R2

En la interfaz S0/0/0 del router 2 (R2), configure la dirección IPv6 que se encuentra en la tabla de direccionamiento, habilite el enrutamiento dinámico de OSPFv3 en la interfaz, utilizando los siguientes comando:

```
R2(config)# interface s0/0/0
R2(config-if)# ipv6 address 2001:4:4:4::2/64
R2(config-if)# clock rate 64000
R2(config-if)# ipv6 ospf 1 area 0
R2(config-if)# no shutdown
R2(config-if)# end
```



Configure la dirección IP que se encuentra en la tabla de direccionamiento de la interfaz Serial0/0/1 en (R2)

```
R2(config)# interface s0/0/1  
R2(config-if)# ipv6 address 2001:5:5:5::1/64  
R2(config-if)# clock rate 64000  
R2(config-if)# ipv6 ospf 1 area 0  
R2(config-if)# no shutdown  
R2(config-if)# end
```

Para configurar OSPF en el router se utilizan los siguientes comandos

```
R2(config)# ipv6 router ospf 1  
R2(config-rtr)# router-id 10.2.2.2
```

Configurar la dirección IPv6 y OSPFv3 en R3

En la interfaz S0/0/0 del router 3 (R3), configure la dirección IPv6 que se encuentra en la tabla de direccionamiento, utilizando los siguientes comandos:

```
R3(config)# interface s0/0/0  
R3(config-if)# ipv6 address 2001:6:6:6::2/64  
R3(config-if)# ipv6 ospf 1 area 0  
R3(config-if)# no shutdown  
R3(config-if)# end
```

Configure la dirección ip que se encuentra en la tabla de direccionamiento de la interfaz Serial0/0/0 en R3.

```
R3(config)# interface s0/0/1  
R3(config-if)# ipv6 address 2001:5:5:5::2/64  
R3(config-if)# ipv6 ospf 1 area 0  
R3(config-if)# no shutdown  
R3(config-if)# end
```

Configure la dirección ip que se encuentra en la tabla de direccionamiento de la interfaz FastEthernet0/0 en R3.

```
R3(config)# interface fa0/0  
R3(config-if)# ipv6 address 2001:3:3:3::1/64  
R3(config-if)# ipv6 ospf 1 area 0  
R3(config-if)# no shutdown  
R3(config-if)# end
```

Emita el comando show IPv6 ospf database para verificar los aspectos específicos de la tabla de topología OSPFv3.

```
Router>en  
Router#show ipv6 route ospf
```



10. Integrated IS-ISv6

IS-IS es un protocolo de la capa de red. Permite a sistemas intermedios (IS's) dentro de un mismo dominio cambiar su configuración e información de ruteo para facilitar la información de encaminamiento y funciones de transmisión de la capa de red.

El protocolo de encaminamiento IS-IS está pensado para soportar encaminamiento en grandes dominios consistentes en combinaciones de muchos tipos de subredes. Esto incluye enlaces punto a punto, enlaces multipunto, subredes X.25 y subredes Broadcast tales como las ISO 8802 LANs. Para poder soportar dominios grandes, la previsión está hecha para que el ruteo intradominio sea organizado jerárquicamente. Un dominio grande puede ser dividido administrativamente en áreas. Cada sistema reside en exactamente un área.

Existen similitudes entre OSPF e IS-IS:

Ambos mantienen una base de datos de los estados de los enlaces (Link state), desde la cual el algoritmo de SPF computa el árbol de SPF.

- Ambos envían Hello para mantener adyacencias.
- Ambos usan áreas para generar topologías jerárquicas.
- Ambos suman entre áreas.
- Ambos son de tipo classless.
- Ambos seleccionan un router designado en redes de tipo Broadcast.
- Ambos tienen capacidad de autenticación.

En IS-IS el equivalente al LSA es el Link State PDU, o LSP, aunque mientras que el LSA tiene el header de OSPF y el header de IP, el LSP es un paquete en sí. En IS-IS existen dos niveles (levels), los cuales hacen las funciones de áreas. Un router L1 tendrá conocimiento solo de su misma área. Un router L2 sabrá solamente rutas de backbone. Y L1/L2 que son como los ABR ya que conocen de ambas topologías. Los routers cisco por default son L1/L2.

Entonces, un router de tipo L1 no tiene conocimiento alguno de las redes que se encuentran fuera de su área, a diferencia de OSPF donde los ABR por default envían las rutas Interárea dentro de las áreas. Este comportamiento de IS-IS es más parecido a un área de tipo totally stub. Donde el router interno envía todo lo que no conoce a los routers de nivel 2.

El router de nivel 2 envía dentro de los LSP con el bit Attached encendido ATT que dice que él puede alcanzar a otras áreas. Network entity titles NET.



Aunque IS-IS puede ser implementado en una red de tipo IP, aún sigue siendo el protocolo de CLNP de ISO, es por ello que es necesario que tenga una dirección de tipo ISO ya que sus PDU aún son de este stack; a estas direcciones se les llama NET.

La NET se compone del Área ID (longitud variable), System ID (6 octetos) y selector (1 octeto).



El área ID es el identificador del área, y es un número variable. El System ID es un número identificador del router o ES, y debe de ser único, generalmente se asigna la MAC Address del equipo. Finalmente el “Sel” es un número que como TCP/IP identifica el tipo de servicio de capa de transporte, siempre en cuestiones de ruteo este valor será 0.

10.1 Vecinos y Adyacencias

Los Hello se envían cada 10 segundos en redes de tipo Broadcast, a diferencia de OSPF este intervalo no debe ser el mismo para que se forme una adyacencia.

Para que se forme una adyacencia se deben de cumplir una de las siguientes condiciones:

- Dos routers L1 con el mismo Área ID.
- Dos routers L2 sin importar diferente Área ID.
- Router L1 con L1/L2 si es misma Área ID.
- Router L2 con L1/L2 sin importar diferente Área ID.

Al igual que en OSPF los Hello sirven como Keepalives, y en ellos se incluyen el hold time que es de 3 veces el Hello. La tabla de vecinos de IS-IS se puede ver con el comando:

```
R1#sh cns is-neighbors
System Id      Interface  State  Type Priority Circuit Id  Format
R2             Et0/0     Up     L1L2 64/64   R2.01      Phase V
```

El estado puede ser “INIT o UP”.

La prioridad es usada para seleccionar al DR. Podemos observar que existen 2 números de 64. Este valor es el valor por default. Son dos porque uno es para la elección de la adyacencia de L1 y el otro para el L2. Si existe empate entre las prioridades de la interfaz, el desempate es quien tenga el MAC Address más alto. En los enlaces de tipo punto a punto o redes non Broadcast donde no se elige DR la prioridad será cero.



El circuit ID es un identificador de la interfaz, cuando esta interfaz se encuentra en una red de tipo Broadcast el System ID del DR se concatena con el Pseudonode (se explica más adelante) y se llama LAN ID.

10.2 Procesos de ISIS – Update Process

Este proceso es el responsable de construir las bases de datos de L1 y L2. Para lograrlo los paquetes de L1 se inundan en toda el área, mientras que los de L2 se mandan por todas las adyacencias de L2.

Max Age

Cada LSP tiene un contador que va decrementando desde MaxAge (20 minutos) hasta cero. Al igual que OSPF en ISIS la tabla se tiene que refrescar cada cierto tiempo y evitar que llegue el conteo a cero.

La ventaja de que ISIS este contador vaya en decremento, es que se puede aumentar de 20 minutos hasta 65535 que son 18.2 horas. Esto en ambientes estables evita la retransmisión de la base de datos. Este valor se cambia con max-lsp-lifetime, y para cambiar el intervalo de refresco de las rutas es con el comando lsp-refresh-interval.

Checksum

También dentro de este proceso se tiene un Checksum, si llega un LSP con un Checksum incorrecto, el router lo descarta, en los router cisco el comando ignore-lsp-errors se encuentra habilitado, de lo contrario si el router recibiera un LSP incorrecto lo que haría sería poner el maxAge en cero y reenviarlo al router que envió el paquete.

Sequence number

Es un numero de 32 bits que va incrementado de 1 en 1, cuando llega al máximo el router dueño de esos LSP apaga el proceso de ISIS durante 21 minutos (Max-Age+ZeroAgeLifetime), esto con el fin de que todos sus LSP se purguen de las BD de sus vecinos y comience otra vez la cuenta en cero.

MAC Address de Multicast.

En redes de tipo punto a punto los LSP de L1 y de L2 se envían directamente al vecino. En las redes Broadcast, las LSP se envían con la MAC destino Multicast de:

- ❖ AIL1ISs – 0180.c200.0014
- ❖ AIL2ISs – 0180.c200.0015

SNP



ISIS utiliza los SNP para ACK de LSP, request y mantener la sincronización de la BD. Existen dos SNP los Partial SNP y los Complete SNP. Los CNSP son como los DBD de OSPF, y los PNSP son usados como request y ACK.

```
R1#sh isis database

IS-IS Level-1 Link State Database:

LSPID          LSP Seq Num  LSP Checksum  LSP Holdtime  ATT/P/OL

R1.00-00      * 0x000000CD  0x80CF        464           0/0/0

R2.00-00      0x000000CD  0x4F78        716           0/0/0

R2.01-00      0x000000CD  0x8CFC        691           0/0/0

IS-IS Level-2 Link State Database:

LSPID          LSP Seq Num  LSP Checksum  LSP Holdtime  ATT/P/OL

R1.00-00      * 0x000000CC  0x240D        744           0/0/0

R2.00-00      0x000000CB  0x15B4        859           0/0/0

R2.01-00      0x000000CB  0x20F3        414           0/0/0
```

El asterisco significa que dicho LSP es generado por el mismo router.

Cada DIS (Designated router) genera un Pseudonode ID por nivel uno para L1 y otro para L2, que es el que representa al nodo Broadcast (similar al network type 2 de LSA). En la salida anterior marcado de amarillo podemos ver el LSP enviado por el router y el Pseudonode enviado en representación del nodo.

El un mero marcado en rojo es el Pseudonode ID que es el SysID del router DIS + un numero identificador (nunca será cero ya que cero es para el mismo router), que es el generado por el DIS en una red de tipo Broadcast, el Azul es un numero identificador de cuando el LSP es superior al MTU soportado por la Interfaz entonces este LSP se fragmenta, este LSP constara del mismo System ID + mismo Pseudonode + diferente identificador.

Procesos de ISIS – Decision Process

Una vez que el proceso de Update ha llenado la BD, se corren dos procesos SPF uno por L1 y otro por L2 para seleccionar las mejores rutas.

La métrica en los routers cisco es la suma del valor asignado (0 a 63) de las interfaces de salida. Por default todas las interfaces de salida tienen un valor de 10 para las L1 y L2. Este valor se puede cambiar con el comando

```
(config-if)#isis metric # //Valor de 0 a 63.
```



También en ISIS existen las rutas internas y externas. Las internas son de ISIS ya sean L1 o L2 y las externas son de otros IGP.

Una ruta L1 es preferida sobre una L2, después la decisión es el costo de la ruta. ISIS ingresa a la tabla de ruteo de hasta 6 rutas de igual costo para realizar el balanceo.

Cuando un router de L1 intenta mandar sus paquetes hacia fuera de su área, este no conoce las rutas hacia afuera, entonces en ISIS los routers de “borde” mandan en el LSP un bit llamado ATT el cual si esta encendido indica que ellos pueden llegar a otras áreas, los routers de L1 seleccionan al router con el ATT encendido que mejor métrica tenga, por ejemplo se puede observar que R1 se seleccionó como el mejor router para alcanzar a las otras áreas:

```
R3#sh isis database

IS-IS Level-1 Link State Database:

LSPID          LSP Seq Num  LSP Checksum  LSP Holdtime  ATT/P/OL
-----
R1.00-00       0x000000D2   0x5EC4        202           1/0/0
R2.00-00       0x000000D4   0x3404        303           0/0/0
R2.01-00       0x000000D2   0x8202        440           0/0/0
R3.00-00       * 0x0000000A 0x7017        655           0/0/0

IS-IS Level-2 Link State Database:

LSPID          LSP Seq Num  LSP Checksum  LSP Holdtime  ATT/P/OL
-----
R3.00-00       * 0x0000000A 0x2860        654           0/0/0
```

Este mecanismo de ATT no es siempre el más eficiente. Ya que si tenemos dos routers con el ATT encendido que están enviando la ruta por default a un L1, este router de L1 seleccionara como salida de los paquetes a uno de los routers de borde, siendo este el que le parezca con mejor métrica.

IS-ISv6 prevé la inclusión de campos de longitud variable (TLV) en todo los paquetes IS-IS (Hello, LSP, y SNP). La Información de direccionamiento relevante se almacena en campos TLV. Los paquetes Hello y paquetes LSP llevan un campo que especifica los protocolos de capa de red. Cada protocolo de capa de red con soporte, se especifica por su NLPID, asignado por la ISO. El valor de IPv6 NLPID es 142 (0x8E).

10.3 La Capacidad IPv6 TLV (Tipo 236)

La accesibilidad TLV de IPv6 en IS-ISv6 corresponde directamente a la accesibilidad común TLV (tipo-longitud-valor) y extendida accesibilidad TLV en IS-IS. La siguiente figura muestra el formato IPv6 TLV de accesibilidad.

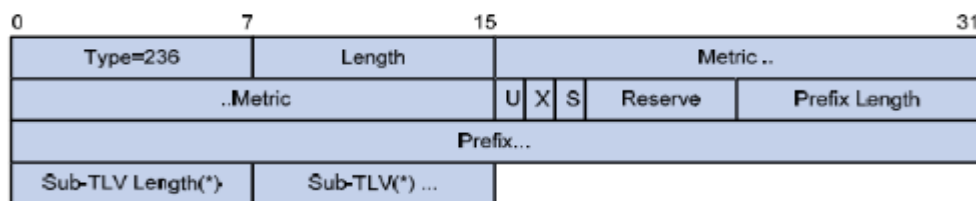


Figura 84: IPv6 reachability TLV format.

Los campos son:

- Tipo: Un valor de 236 indica que este TLV es un TLV de accesibilidad IPv6.
- Longitud: longitud TLV.
- Métricas: extensión métrica, en el rango de 0 a 4261412864. Si el valor de la métrica es mayor que 4261412864, la información de accesibilidad IPv6 será ignorado.
- T: Bit de activado o desactivado que se utiliza para evitar los bucles de enrutamiento. Cuando una ruta se anuncia desde un router de nivel-2 a uno de nivel-1, este campo se pone a 1 para evitar la ruta de ser devuelta en bucle.
- X: Bit de Ruta de redistribución. Un valor de 1 significa que la ruta se redistribuye desde otro protocolo.
- S: Si un TLV no lleva ningún sub-TLV, este campo se establece en 0; de lo contrario, se establece en 1, es decir, el prefijo IPv6 es seguido por la información sub-TLV.
- Reserve: Este campo está reservado.
- Prefix Length: Longitud de prefijo de ruta IPv6.
- Prefijo: prefijo de ruta IPv6.
- Sub-TLV / Sub-TLV Longitud: Sub-TLV y su longitud, opcional.

Interfaz de Dirección IPv6 TLV (Tipo 232)

Mapas de direcciones IPv6 TLV de la interfaz directamente a la dirección de la interfaz TLV IPv4 en IS-IS.

La siguiente figura muestra el formato IPv6 dirección de la interfaz TLV:

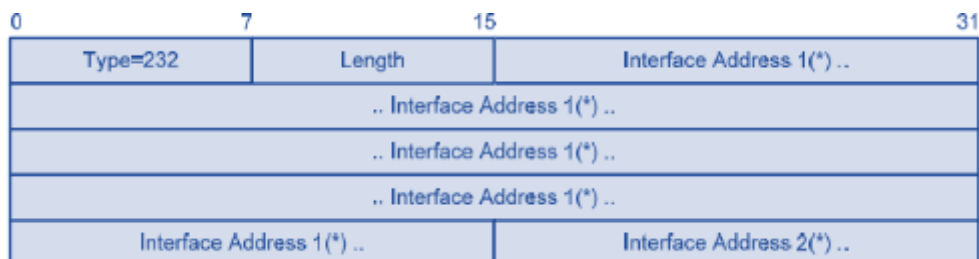


Figura 85: IPv6 interface address TLV format.

Los campos significativos se describen a continuación:

- Tipo: Un valor de 232 indica que este TLV es una interfaz TLV de dirección IPv6.



- Longitud: longitud TLV.
- Interfaz Dirección: dirección IPv6 de la interfaz. Para PDUs hello de las "Dirección de Interfaces" TLV debe contener sólo las direcciones IPv6 locales de vínculo asignadas a la interfaz que está enviando el paquete Hello. Para LSP, las "Interfaces Dirección" TLV debe contener solamente la no-enlace-local(non-link-local) direcciones de IPv6, asignadas a la SI, a saber, las direcciones de unidifusión global IPv6 asignada a la interfaz.
- *: Significa esto es opcional.

Protocolo de TLV

Un protocolo TLV especifica el protocolo soportado en el router original. El protocolo TLV tiene un NLPID correspondiente a cada protocolo soportado. Para un router de apoyo IS-ISv6, un campo NLPID con un valor de 0x81 necesita ser añadido a la TLV.

Type=129
Length
NLPID
:
NLPID

Figura 86: Protocol TLV.



10.4 IS-ISv6 Adyacencia

IS-IS utiliza paquetes de saludo para descubrir routers vecinos y establecer adyacencias con ellos. Después de establecer una adyacencia entre dos routers, envían periódicamente Hola paquetes entre sí para mantener la adyacencia. En IS-ISv6, el paquete de saludo es extendido para soportar IPv6:

- Un NLPID 8 bits se añade al protocolo de TLV. Un valor de 0x81 para este NLPID significa que los soportes actuales del router IS-ISv6.
- En los paquetes de saludo, se añade la dirección de la interfaz IPv6 TLV y está lleno de la dirección de enlace local IPv6.

Las características de IPv6 para IS-IS también permiten que se sumen a las rutas IPv4, los prefijos IPv6. Se crea un nuevo address family para incluir IPv6. IS-IS IPv6 soporta tanto single-topology como multiple topology.

10.5 IS-IS Single Topology:

- IS-IS tiene la particularidad de soportar múltiples protocolos de capa 3.
- Si tenemos IS-IS con otro protocolo (por ej.: IPv4) configurado en una interfaz, podemos configurar también ISIS Para IPv6.
- Todas las interfaces deben ser configuradas en forma idéntica en cada address family (misma topología), tanto para los routers L1 como los L2.

10.6 IS-IS multi-topology

- Permite mantener topologías independientes dentro de un área.
- Elimina la restricción para todas las interfaces de tener idénticas topologías por cada Address familiy.
- Los routers construyen una topología por cada protocolo de capa 3, por lo que, pueden encontrar el camino optimo (SPF) aun si algún link soporta solo uno de estos protocolos.



10.7 Configuración Integrated IS-IS

Configurar protocolo IS-IS

```
Router> enable
Router# configure terminal
Router(config)# router isis <area-name>
Router(config-router)# net <network-entity-title>
```

Configurando la interfaz

```
Router> enable
Router# configure terminal
Router(config)# interface <type> <number>
Router(config-if)# ipv6 address <ipv6-prefix/prefixlength>
Router(config-if)# ipv6 router isis <area-name>
```

IS-IS mult topology

```
Router> enable
Router# configure terminal
Router(config)# router isis <area-name>
```



11. BGP4

Border Gateway Protocol versión 4 (BGP) el protocolo ha existido desde 1994 y fue actualizado muchas veces estos últimos años. BGP4 está definido en el estándar RFC 4271.

Es un protocolo de encaminamiento usado entre sistemas autónomos que poseen internet. BGP Externo (EBGP) se usa entre sistemas autónomos y el BGP Interno (IBGP) es usado dentro de un sistema autónomo.

Se basa en el PVP (Path Vector Protocol).

- Similar al de Distancia Vector
- Cada encaminador frontera envía a sus vecinos ("peerings") la ruta completa a un destino, no solo la distancia.
- El camino (path) es una secuencia de ASs hasta el destino.

Al pasar los años, BGP se extendió para llevar diferentes tipos de información de enrutamiento. En RFC 4760, existen muchas extensiones multiprocesos para BGP-4, permite que BGP opere sobre IPv4 o IPv6 y transporte cualquier tipo de información de enrutamiento.

BGP es el sistema nervioso central con el que prácticamente todos los proveedores de servicios están conectados. Debido a que BGP es el protocolo de enrutamiento que se usa en Internet, es objeto de ataques. Los atacantes saben que si pueden encontrar una debilidad en BGP y explotarla, podrían potencialmente desestabilizar toda la Internet. En RFC 4272, "Análisis de Vulnerabilidades de Seguridad BGP," mostró las debilidades en BGP que los proveedores de servicios deben tratar de evitar. Por lo tanto, es importante que trabaje para asegurar BGP, centrándose en las siguientes áreas:

- Autenticación.
- Confidencialidad.
- Integridad.
- Disponibilidad.

Convencionalmente hay varios enfoques de seguridad de sesiones BGP, incluyendo el siguiente:

- Explícitamente configuración pares BGP
- El uso de sesión BGP de claves compartidas
- Aprovechamiento de un túnel IPsec
- El uso de direcciones loopback en pares BGP
- El control de los paquetes BGP Time-to-Live (TTL)
- Filtrado en interfaces pares.
- Uso de pares de enlace local (peering link-local)



- La prevención de rutas largas de AS.
- Limitar el número de prefijos recibidos.
- Prevención de actualizaciones BGP que contienen números privados de AS.
- Maximización de la disponibilidad de pares BGP.
- Registro de actividad del vecino par (peer).
- Asegurar el IGP
- Medidas extremas para asegurar las comunicaciones entre los BGP.

11.1 Configuración Explícita Peers (pares) BGP

Una técnica para asegurar sesiones BGP es el concepto de que las sesiones BGP deben ser configuradas en cada router contiguo. Los acuerdos de interconexión se realizan de forma explícita por ambos comunicándose vía BGP.

Por lo tanto, un router no establecerá una sesión de intercambio con otro router que no ha sido configurado previamente con su par, y ambos estarán de acuerdo entre sí sobre la configuración de BGP. Una sesión BGP no se establece si sólo hay un router configurado. Debe haber configuraciones complementarias en ambos para que se puedan comunicar. Las comunicaciones en BGP se realizan a través de TCP, por lo que el protocolo debe confiar en un formato debidamente configurado sobre la capa IP (IP-Layer).

BGP utiliza el puerto TCP 179, por lo que tiene algo de seguridad inherente al hecho que es un protocolo orientado a la conexión. Se mantiene el estado de sesión TCP entre ambos.

El hecho de que BGP es un protocolo de enrutamiento de estado de capa de transporte normalmente proporciona cierto nivel de seguridad, pero también es una de las debilidades de BGP. Los atacantes pueden suplantar paquetes BGP y enviarlos hacia uno de los routers BGP, o podría atacar la sesión TCP entre dos routers BGP. Amenazas contra sesiones largas de TCP, implican secuestro de sesiones TCP donde se utiliza el número de secuencia, y así tomar el control de uno de los pares BGP. Una solución a este problema es implementar un número de secuencia aleatorio. Por lo tanto adivinar el siguiente número de secuencia o acuse de recibo (ACK) sería difícil e improbable.

11.2 Uso de Claves Compartidas en Sesiones BGP

Uno de los métodos más utilizados de asegurar las comunicaciones BGP es el uso compartido secreto (contraseña). RFC 2385, "La protección de Sesiones BGP vía TCP usando la Firma MD5", define como una simple contraseña se puede utilizar con un mensaje que usa algoritmo 5 (MD5) insertándolos en los paquetes BGP. Esto añade autenticación de BGP y ayuda a evitar que un atacante spoofing (envenene) un par BGP.



A pesar de que es una buena práctica utilizar una contraseña diferente para cada intercambio de sesión, esta puede ser difícil de mantener. No obstante, es prudente utilizar la misma clave secreta para todas las sesiones peering (pares). Como se suele decir, no es un secreto si le dices a un montón de gente. RFC 3562, "Claves Consideradas de gestión para TCP con Firma MD5", se define como un sistema centralizado puede mantener la seguridad de las claves para todas las organizaciones. En un router Cisco, la contraseña se le asigna en el momento de que el vecino es configurado. A continuación se presenta la configuración de un router usando comando para habilitar la autenticación MD5 para un par BGP:

```
neighbor neighbor-ipv6-address password P@ssw0rd
```

11.3 Aprovechamiento de un Túnel IPsec

Otra técnica para asegurar las comunicaciones BGP es aprovechar la seguridad de un túnel IPsec. IPsec es una manera fuerte para asegurar pares BGP, proteger la integridad de las actualizaciones, y ayudar en la prevención de ataques de denegación de servicio que se dirigen a los pares BGP. El uso de IPsec es mejor que MD5 debido a que se mantiene actualizando claves a cada momento. Debido a que BGP es un protocolo TCP, puede utilizar IPsec sin modificación alguna. Sin embargo, una conexión IPsec debe ser creada para su interconexión.

Esto puede agregar significativamente una sobrecarga a los routers, por lo que podría consumir muchos recursos en términos de CPU. Configuración y solución de problemas del túnel IPsec puede añadir una carga significativa para el mantenimiento de una red de proveedor de servicios. Además, el túnel IPsec que se utiliza para el envío de información de enrutamiento se utiliza así para reenviar tráfico. La sobrecarga de paquetes de tamaño añadido que IPsec añade impactaría negativamente en el rendimiento de procesamiento. A pesar de que el uso de IPsec es un método seguro, no se utiliza ampliamente.

Aun así, un atacante que sabe que un router está usando autenticación, simplemente puede crear un gran número de paquetes falsificados con parámetros de autenticación y enviarlos hacia ese router. Esto haría que el router para procesar estos paquetes falsos (incluso si han sido rechazados de forma rápida y artificialmente) consuma muchos recursos de ese router. De esa manera el proceso de verificación de paquetes se puede retrasar, y así el atacante podría cumplir con su objetivo. Los atacantes podrían lanzar muchos errores de autenticación en el router BGP para potencialmente irrumpir en él. Por lo tanto, la autenticación no puede ser el único método de asegurar las comunicaciones BGP.



11.4 Configuración de BGP4

Ejemplo:

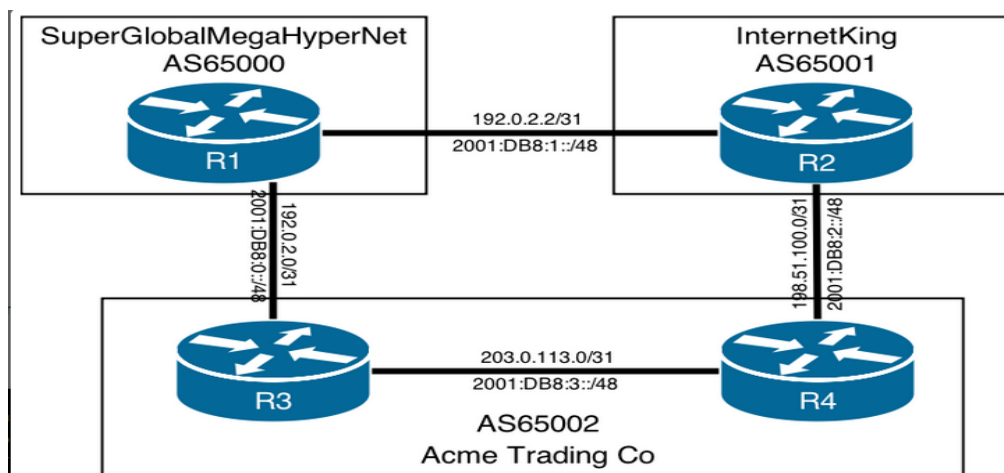


Figura 87: Topología BGP

Primero hacemos la conectividad básica de IPv6:

R1

```
R1#sh ipv6 int bri
GigabitEthernet0/0      [up/up]
    FE80::C801:2BFF:FE86:8
    2001:DB8::1
GigabitEthernet1/0      [up/up]
    FE80::C801:2BFF:FE86:1C
    2001:DB8:1::1
```

R2:

```
R2#sh ipv6 int bri
GigabitEthernet0/0      [up/up]
    FE80::C802:2BFF:FE86:8
    2001:DB8:2::1
GigabitEthernet1/0      [up/up]
    FE80::C802:2BFF:FE86:1C
    2001:DB8:1::2
```



R3:

```
R3#sh ipv6 int bri
GigabitEthernet0/0      [up/up]
    FE80::C803:2BFF:FE86:8
    2001:DB8::2
GigabitEthernet1/0      [up/up]
    FE80::C803:2BFF:FE86:1C
    2001:DB8:3::2
```

R4:

```
R4#sh ipv6 int bri
GigabitEthernet0/0      [up/up]
    FE80::C800:2BFF:FE86:8
    2001:DB8:2::2
GigabitEthernet1/0      [up/up]
    FE80::C800:2BFF:FE86:1C
    2001:DB8:3::2
```

Ahora veremos todas las direcciones almacenada y para comprobar su funcionalidad se hará ping

Ping R2 para R1:

```
R1#ping 2001:DB8:1::2
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 2001:DB8:1::2, timeout is 2
seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 2
0/24/40 ms
```

Ping R3 para R1

```
R1#ping 2001:DB8::2
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 2001:DB8::2, timeout is 2 s
econds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 2
0/27/44 ms
```



Ping R4 para R3

```
R3#ping 2001:DB8:3::2
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 2001:DB8:3::2, timeout is 2
seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 2
0/24/32 ms
```

Ping R2 para R4

```
R4#ping 2001:DB8:2::1
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 2001:DB8:2::1, timeout is 2
seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max = 2
0/25/36 ms
```

Configuración de IPv6 BGP

R1:

```
R1#conf t
R1(config)#ipv6 unicast-routing
R1(config)#router bgp 65000
R1(config-router)#address-family ipv6
R1(config-router-af)#redistribute connected
R1(config-router-af)#neighbor 2001:DB8:1::2 remote-as 65001
R1(config-router-af)#neighbor 2001:DB8::2 remote-as 65002
```

R2:

```
R2#conf t
R2(config)#ipv6 unicast-routing
R2(config)#router bgp 65001
R2(config-router)#address-family ipv6
R2(config-router-af)#redistribute connected
R2(config-router-af)#neighbor 2001:DB8:1::1 remote-as 65000
R2(config-router-af)#neighbor 2001:DB8:2::2 remote-as 65002
```



R3:

```
R3#conf t
R3(config)#ipv6 unicast-routing
R3(config)#router bgp 65002
R3(config-router)#address-family ipv6
R3(config-router-af)#redistribute connected
R3(config-router-af)#neighbor 2001:DB8::1 remote-as 65000
R3(config-router-af)#neighbor 2001:DB8:3::2 remote-as 65002
```

R4:

```
R4#conf t
R4(config)#ipv6 unicast-routing
R4(config)#router bgp 65002
R4(config-router)#address-family ipv6
R4(config-router-af)#redistribute connected
R4(config-router-af)#neighbor 2001:DB8:2::1 remote-as 65001
R4(config-router-af)#neighbor 2001:DB8:3::1 remote-as 65002
```

Verificación BGP IPv6

R1:

```
R1#sh ip bgp ipv6 unicast summary
BGP router identifier 192.0.2.2, local AS number 65000
BGP table version is 7, main routing table version 7
4 network entries using 596 bytes of memory
8 path entries using 608 bytes of memory
10/3 BGP path/bestpath attribute entries using 1240 bytes of memory
3 BGP AS-PATH entries using 72 bytes of memory
1 BGP community entries using 24 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
BGP using 2540 total bytes of memory
BGP activity 15/0 prefixes, 25/1 paths, scan interval 60 secs
Neighbor      V    AS MsgRcvd MsgSent  TblVer  InQ OutQ Up
/Down  State/PfxRcd
2001:DB8::2    4 65002    13     14      7    0   0 00
:05:32        3
2001:DB8:1::2  4 65001    16     17      7    0   0 00
:08:03        3
R1#!Lets ping R4 from here, it isn't directly connected so ha
s
```



R2:

```
R2#sh ip bgp ipv6 unicast summary
BGP router identifier 198.51.100.0, local AS number 65001
BGP table version is 6, main routing table version 6
4 network entries using 596 bytes of memory
8 path entries using 608 bytes of memory
10/3 BGP path/bestpath attribute entries using 1240 bytes of
memory
3 BGP AS-PATH entries using 72 bytes of memory
2 BGP community entries using 48 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
BGP using 2564 total bytes of memory
BGP activity 12/0 prefixes, 22/1 paths, scan interval 60 secs
Neighbor      V    AS MsgRcvd MsgSent  TblVer  InQ OutQ Up
/Down  State/PfxRcd
2001:DB8:1::1  4 65000      17      16        6    0    0 00
:08:05        3
2001:DB8:2::2  4 65002       8      10        6    0    0 00
:02:35        3
R2#!Lets ping R3 from here, that will have to
R2#!go via R1 or R4
```



12. VLANs (Red de área local virtual)

Es una manera de independizar redes lógicamente dentro de una misma red. Varias VLAN pueden coexistir en un único conmutador físico o en una única red física. Son útiles para reducir el tamaño del dominio de difusión y ayudan en la administración de la red, separando segmentos lógicos de una red de área local.

Las VLANs se encuentran conformadas por un conjunto de dispositivos de red interconectados (hubs, bridges, switches), la definimos como una subred por software y es considerada como un dominio de broadcast que pueden estar en el mismo medio físico o bien pueden estar sus integrantes ubicados en distintos sectores de la corporación.

La tecnología de las VLANs se basa en el empleo de switches, en lugar de hubs, de tal manera que esto permite un control más inteligente del tráfico de la red.

Las consideraciones de VLANs para IPv6 son los mismos que para IPv4. Cuando se utilizan las configuraciones de doble pila (DUAL-STACK), IPv6 e IPv4 atraviesan la misma VLAN. Cuando se utiliza un túnel, el tráfico IPv6 e IPv4 del túnel atraviesa la VLAN.

12.1 Tipos de VLANs

VLANs de puerto central

Es en la que todos los nodos de una VLAN se conectan al mismo puerto del switch.

VLANs estáticas

Las asignaciones en una VLAN estática se crean mediante la asignación de los puertos de un switch o conmutador a dicha VLAN. Cuando un dispositivo entra en la red, automáticamente asume su pertenencia a la VLAN a la que ha sido asignado el puerto. Si el usuario cambia de puerto de entrada y necesita acceder a la misma VLAN, el administrador de la red debe cambiar manualmente la asignación a la VLAN del nuevo puerto de conexión en el switch.

Los puertos de los switches están ya pre-asignados a las estaciones de trabajo y poseen las siguientes características:

- Por puerto
- Por dirección MAC
- Por protocolo
- Por direcciones IP
- Por nombre de usuarios



VLANS dinámicas

Las VLAN dinámicas son puertos del switch los que automáticamente determinan a que VLAN pertenece cada puesto de trabajo. El funcionamiento de estas VLANs se basa en las direcciones MAC, direcciones lógicas o protocolos utilizados. Cuando un puesto de trabajo pide autorización para conectarse a la VLAN, el switch revisa la dirección MAC ingresada previamente por el administrador en la base de datos de las mismas, y automáticamente se configura el puerto al cual corresponde por la configuración de la VLAN. El mayor beneficio de las VLAN, es el menor trabajo de administración dentro del armario de comunicaciones cuando se cambian de lugar las estaciones de trabajo o se agregan, y también notificación centralizada cuando un usuario desconocido pretende ingresar a la red.

12.2 Tipos de puertos

Un switch que utiliza VLANs puede tener dos tipos de puertos: puertos de acceso y puertos de Trunk. A continuación se da una explicación de cada uno de ellos.

- **Puertos de acceso:** Este tipo de puertos son los que conectan hosts finales. Trabajan con las tramas clásicas de Ethernet, sin el agregado de las etiquetas de VLAN.
- **Puertos de Trunk:** Los puertos de Trunk tienen una función especial que es la de conectar switches entre sí o un switch con un router. Cuando llega tráfico a un puerto de Trunk proveniente desde el propio switch, éste es etiquetado con el identificador de VLAN y enviado por el puerto. El equipo que lo recibe, desencapsula la trama Ethernet (quitándole la etiqueta) y lo envía al puerto que corresponda.

12.3 Interconexiones de switch con VLANs y puertos trunk

Este puerto tiene las características de pertenecer a diferentes o a todas las VLANs; es decir, solo se necesita un solo cable para conectar a todas la VLANs, transporta el tráfico de cada una de ellas. Un enlace troncal puede conectar:

- Un switch a otro switch
- Un switch a un router
- Un switch a un servidor instalando una NIC especial que admite enlace troncal.

Para conseguir conectividad entre VLAN a través de un enlace troncal entre Switchs, las VLAN deben estar configuradas en cada switch.

El VLAN Trunking Protocol (VTP) proporciona un medio sencillo de mantener una configuración de VLAN coherente a través de toda la red conmutada. VTP permite soluciones de red conmutada fácilmente escalable a otras dimensiones, reduciendo la necesidad de configuración manual de la red.



La siguiente figura está formada por un router y 2 switches, se implementará un direccionamiento IPv6. Todas las redes que están presentes en el router usan un prefijo /64. Se crearán 3 VLANs con sus respectivos nombres y demás configuraciones.

12.4 Configuración de VLANs con IPv6.

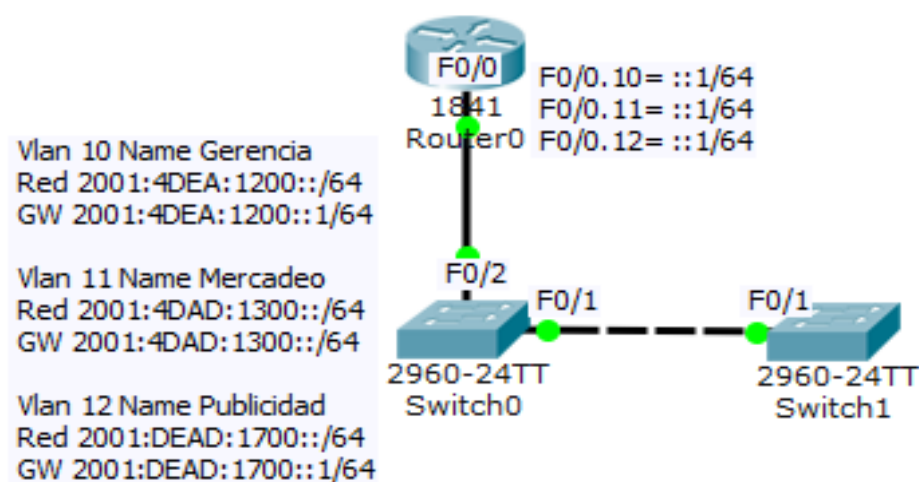


Figura 88: Topología de VLANs.

Lo primero que se debe de hacer es indicar que se usará el direccionamiento IPv6.

```
Router>enable
Router#config t
Enter configuration commands, one per line. End with CNTL/Z.
Router(config)#ipv6 unicast-routing
Router(config)#
```

Configuración de cada VLANs en los respectivos switches

Switch0

```
Switch>enable
Switch#vlan database
% Warning: It is recommended to configure VLAN from config mode,
as VLAN database mode is being deprecated. Please consult user
documentation for configuring VTP/VLAN in config mode.

Switch(vlan)#vlan 10 name Gerencia
VLAN 10 added:
  Name: Gerencia
Switch(vlan)#vlan 11 name Mercadeo
VLAN 11 added:
  Name: Mercadeo
Switch(vlan)#vlan 12 name Publicidad
VLAN 12 added:
  Name: Publicidad
Switch(vlan)#ex
```



Asignar las VLANs antes creadas a su interfaz correspondiente y así brindar acceso a ellas.

Switch0

```
Switch>en
Switch#conf t
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)#int f0/1
Switch(config-if)#swi
Switch(config-if)#switchport acc
Switch(config-if)#switchport access vlan
Switch(config-if)#switchport access vlan 12
Switch(config-if)#do w
```

Cambiar el tipo de acceso predefinido de las interfaces de los switches. se deberá de pasar de modo Access a modo troncal.

Switch0

```
Switch#conf t
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)#int f0/2
Switch(config-if)#swi
Switch(config-if)#switchport ac
Switch(config-if)#switchport mo
Switch(config-if)#switchport mode tr
Switch(config-if)#switchport mode trunk |
```

Creación de VTP en los switches tanto en modo servidor como cliente, además se le asignará un dominio y una contraseña de seguridad.

Switch0

```
Switch>en
Switch#conf t
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)#vtp mode server
Device mode already VTP SERVER.
Switch(config)#vtp domain vlanipv6
Changing VTP domain name from NULL to vlanipv6
Switch(config)#vtp password vlan
Setting device VLAN database password to vlan
```

Switch1

```
Switch>en
Switch#conf t
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)#vtp mode client
Setting device to VTP CLIENT mode.
Switch(config)#vtp domain vlanipv6
Changing VTP domain name from NULL to vlanipv6
Switch(config)#vtp password vlan
Setting device VLAN database password to vlan
Switch(config)#do w
```



Creación de sub-interfaces en interfaz 0/0 en el router para que se puedan comunicar las diferentes VLANs existentes conectadas directamente a él.

```
Router>en
Router#conf t
Enter configuration commands, one per line.  End with CNTL/Z.
Router(config)#int f0/0.10
Router(config-subif)#encapsulation dot1q 10
Router(config-subif)#ipv6 address 2001:4DEA:1200::1/64
Router(config-subif)#int f0/0.11
Router(config-subif)#encapsulation dot1q 11
Router(config-subif)#ipv6 address 2001:4DAD:1300::1/64
Router(config-subif)#int f0/0.12
Router(config-subif)#encapsulation dot1q 12
Router(config-subif)#ipv6 address 2001:DEAD:1700::1/64
Router(config-subif)#do w
```

Y por último se debe de activar la interfaz física del router de la siguiente forma.

R0

```
Router>en
Router#conf t
Enter configuration commands, one per line.  End with CNTL/Z.
Router(config)#int f0/0
Router(config-if)#no shutdown
Router(config-if)#
```



13. Frame Relay

Frame Relay es un protocolo de **WAN** que trabaja en las capas físicas y de enlace de datos del modelo de referencia OSI. Nace como sustitutivo del protocolo **X.25**, al ser más eficiente y moderno.

13.1 Beneficios de Frame Relay

- Proporciona un mayor ancho de banda, fiabilidad y elasticidad que las líneas dedicadas.
- Los usuarios de Frame Relay solo pagan por el bucle local (la conexión entre el cliente y el switch Frame Relay del proveedor) y por el ancho de banda contratado. La distancia entre los nodos a comunicar no es importante.
- El ancho de banda es compartido entre muchos usuarios, abaratando los costes.
- Proporciona mayor flexibilidad en el diseño de la red. Ya que, si se quiere conectar un nuevo sitio a un PVC, basta con contratar una línea Frame Relay e indicar el DLCI que tendrá el nuevo extremo y el PVC al que debe conectarse.

13.2 Funcionamiento de Frame Relay

- La conexión entre dos routers extremo es lógica, ya que, internamente en el ISP, no hay un enlace físico que los interconecte, sino que los datos viajan del origen al destino a través de la red Frame Relay.
- Todos los switches Frame Relay por los que viaja un paquete de un origen a un destino, de manera bidireccional, se denomina Circuito Virtual (VC).

La conexión entre un dispositivo DTE y un dispositivo DCE comprende un componente de capa física y un componente de capa de enlace.

El componente físico define las especificaciones mecánicas, eléctricas, funcionales y de procedimiento necesarias para la conexión entre dispositivos. Una de las especificaciones de interfaz de capa física más comúnmente utilizadas es la especificación RS-232.

El componente de capa de enlace define el protocolo que establece la conexión entre los dispositivos DTE, como un router, y el dispositivo DCE, como un switch.

El switch Frame Relay es un dispositivo DCE. Los switches de red mueven tramas desde un DTE en la red y entregan tramas a otros DTE en forma de DCE. Otros equipos informáticos que no se encuentren en la LAN pueden también enviar datos a través de la red Frame Relay.



Dichos equipos utilizan como DTE un dispositivo de acceso Frame Relay (FRAD). A menudo, FRAD hace referencia a un ensamblador/desensamblador de Frame Relay que es un artefacto dedicado o un router configurado para admitir Frame Relay.

Las tramas se mueven de switches a switches a través de la WAN al switch DCE de destino en el extremo de la WAN.

El DCE de destino entrega la trama al DTE de destino.

13.3 Circuitos virtuales

La conexión a través de una red Frame Relay entre dos DTE se denomina circuito virtual (VC, Virtual Circuito). Los circuitos son virtuales dado que no hay una conexión eléctrica directa de extremo a extremo. La conexión es lógica y los datos se mueven de extremo a extremo, sin circuito eléctrico directo. Con los VC, Frame Relay comparte el ancho de banda entre varios usuarios, y cualquier sitio puede comunicarse con otro sin usar varias líneas físicas dedicadas.

Existen 2 formas de establecer circuitos virtuales:

- **Los SVC**, circuitos virtuales conmutados, se definen dinámicamente mediante el envío de mensajes de señalización a la red (CALL SETUP, DATA TRANSFER, IDLE, CALL TERMINATION).
- **Los PVC**, circuitos virtuales permanentes, son pre configurados por la empresa de comunicaciones y, una vez configurados, sólo funcionan en los modos DATA TRANSFER e IDLE. Tenga en cuenta que algunas publicaciones hacen referencia a los PVC como VC privados.

13.4 DLCI (Data Link Connection Identifier)

Es el identificador de canal del circuito establecido en Frame Relay. Este identificador se aloja en la trama e indica el camino a seguir por los datos, es decir, el circuito virtual establecido.

Las conexiones Frame Relay identifican los circuitos virtuales por el identificador de conexión de datos Link (DLCI), estos números asocian una dirección IP con un circuito virtual específico. Los números DLCI sólo tienen importancia local y se asignan generalmente por el proveedor de Frame Relay.

NOTA: Los valores de DLCI tienen importancia local. Lo que significa que solo son únicos para el canal físico en el que residen por lo tanto, los dispositivos de los extremos opuestos de una conexión Pueden usar los mismos valores de DLCI para referirse a diferentes circuitos virtuales.

**DLCI (Data Link Connection Identifier)**

Es el identificador de canal del circuito establecido en Frame Relay. Este identificador se aloja en la trama e indica el camino a seguir por los datos, es decir, el circuito virtual establecido.

Las conexiones Frame Relay identifican los circuitos virtuales por el identificador de conexión de datos Link (DLCI), estos números asocian una dirección IP con un circuito virtual específico. Los números DLCI sólo tienen importancia local y se asignan generalmente por el proveedor de Frame Relay.

NOTA: Los valores de DLCI tienen importancia local. Lo que significa que solo son únicos para el canal físico en el que residen por lo tanto, los dispositivos de los extremos opuestos de una conexión Pueden usar los mismos valores de DLCI para referirse a diferentes circuitos virtuales.

13.5 Mapa Frame Relay

Una tabla en la memoria RAM que define la interfaz remota a la que se asigna un determinado número DLCI. La definición contendrá un número DLCI y un identificador de interfaz siendo típicamente una dirección IP remota. El mapa de Frame Relay puede ser de manera estática o automática según la topología de Frame Relay.

Mapeo dinámico

El router Frame Relay realiza una consulta ARP Inversa en el PVC para descubrir la dirección de capa 3 del siguiente salto y lo almacena en una **tabla de mapeo** (Mapping Table).

Mapeo estático

Consiste en añadir **manualmente** la asociación entre una dirección de capa 3 y el DLCI correspondiente en la tabla de mapeo. Esto es útil cuando se interconectan dispositivos en los que no funciona el ARP Inverso o cuando se trata de una topología **Hub and Spoke**. En este caso concreto, cada nodo (Spoke) no sabe cómo llegar a ningún sitio que no sea el dispositivo que tiene al otro extremo del cable (Hub). Por ello, podría usarse ARP Inverso entre estos dos nodos y realizar un mapeo estático entre el resto de nodos de la topología.

NOTA: No se puede utilizar un mapeo estático y otro dinámico en la misma interfaz. Para ello, podemos deshabilitar el ARP Inverso en la interfaz con el comando `no frame-relay inverse-arp`.



Ejemplo de Mapeo de Circuitos Virtuales con números de puerto

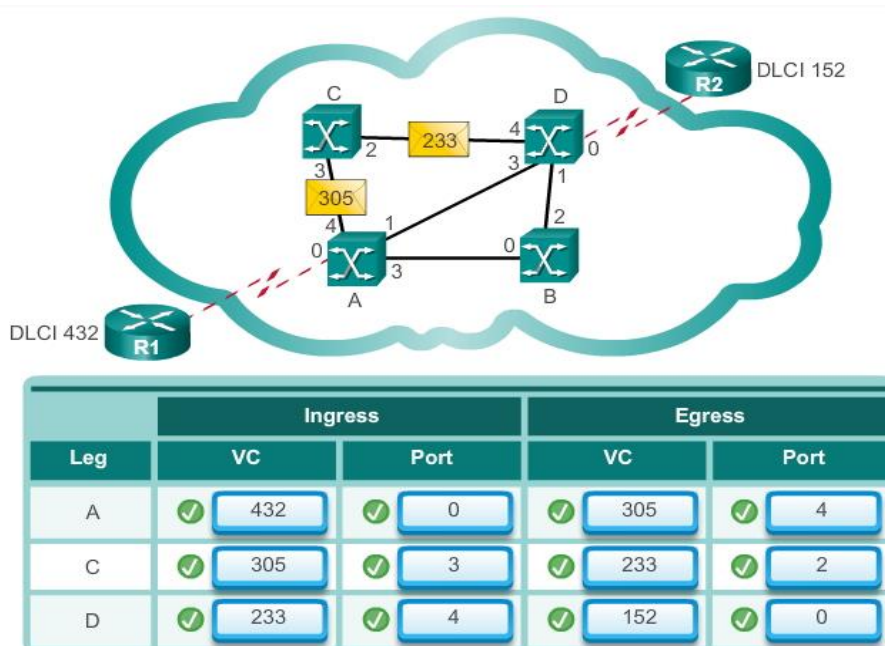


Figura 89: Ejemplo de mapeo de circuitos virtuales.

En este ejemplo tenemos una comunicación que parte de R1 con dirección a R2. Para la pata A (entre R1 y el primer switch Frame Relay) se asocia el DLCI 432 al puerto de acceso (ingress) 0.

Para el siguiente paso, se asocia el DLCI 305 al puerto de salida (egress) 4 del primer switch FR y al puerto de acceso 3 del segundo switch FR.

Para el tercer paso, se asocia el DLCI 233 al puerto de salida 2 del segundo switch FR y al puerto de acceso 4 del cuarto switch FR.

Y, por último, se asocia el DLCI 152 para la pata D al puerto de salida 0 del cuarto switch FR.

ARP inverso

Al contrario que ARP se utiliza para averiguar una dirección MAC en función de una IPv4, Frame Relay utilizar una ARP inversa para conocer la **dirección IPv4** que corresponde con el DLCI del siguiente salto. Esto es necesario antes de que se pueda usar el Circuito Virtual.

ARP Inversa está habilitado por defecto en los dispositivos Cisco y se utilizará solo para los protocolos que estén habilitados en una interfaz.

NOTA: En el caso de IPv6 se utilizan solicitudes y respuestas IND (Inverse Neighbor Discovery)



Subinterfaces

Interfaces virtuales asociados con una interfaz física. Creadas mediante la referencia de la interfaz física, seguido de un período y un número decimal. Para los fines de encaminamiento, sin embargo, sub interfaces son tratados como interfaces físicas. Con sub interfaces, el costo de implementar múltiples circuitos virtuales de Frame Relay se reduce, debido a que sólo un puerto es necesario en el router.

13.6 Encapsulación Frame Relay

Frame Relay toma paquetes de datos de un protocolo de capa de red, como IP o IPX, los encapsula como la parte de datos de una trama Frame Relay y, luego, pasa la trama a la capa física para entregarla en el cable.

Frame Relay acepta un paquete de un protocolo de capa de red como IP. A continuación, lo ajusta con un campo de dirección que incluye el DLCI y una checksum. Se agregan campos señaladores para indicar el comienzo y el fin de la trama. Los campos señaladores marcan el comienzo y el fin de la trama, y siempre son los mismos en números binarios 01111110. Después de haber encapsulado el paquete, Frame Relay pasa la trama a la capa física para su transporte.

El encabezado Frame Relay (campo de dirección) incluye lo siguiente:

- **DLCI.**
- **Dirección extendida (EA).**
- **C/R.**
- **Control de congestión.**

La capa física en general es EIA/TIA-232, 449 ó 530, V.35, o X.21. Las tramas Frame Relay son un subconjunto del tipo de trama HDLC. Por lo tanto, están delimitadas por campos señaladores. El señalador de 1 byte usa el patrón de bits 01111110.

Frame Relay no notifica el origen cuando se descarta una trama.



13.7 Configuración de Frame Relay

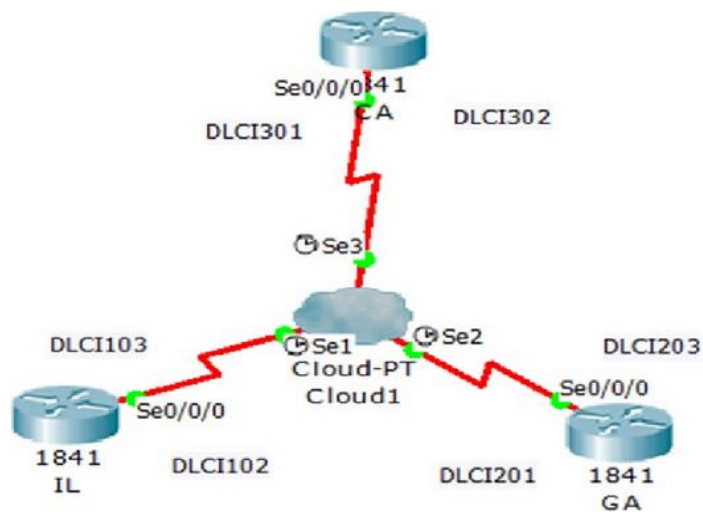


Figura 90: Topología de Frame Relay

Configuración de los routers.

Configuración del router IL

```
int s0/0/0
band 64
encap frame-relay
ipv6 add FC01:CAFE:1:AA11::1/64
ipv6 add FE80::1 link-local
frame-r map ipv6 FC01:CAFE:AA11::3 103
frame-r map ipv6 FC01:CAFE:AA11::2 102
frame-r map ipv6 FE80::3 103
frame-r map ipv6 FE80::2 102
no frame-relay inverse-arp
no shut
```

Configuración del router GA

```
int s0/0/0
band 64
encap frame-relay
ipv6 add FC01:CAFE:1:AA11::2/64
ipv6 add FE80::2 link-local
frame-r map ipv6 FC01:CAFE:AA11::1 201
frame-r map ipv6 FC01:CAFE:AA11::3 203
frame-r map ipv6 FE80::1 201
frame-r map ipv6 FE80::3 203
no frame-relay inverse-arp
no shut
```

Configuración del router CA

```
int s0/0/0
band 64
encap frame-relay
ipv6 add FC01:CAFE:1:AA11::3/64
ipv6 add FE80::3 link-local
frame-r map ipv6 FC01:CAFE:AA11::2 302
frame-r map ipv6 FC01:CAFE:AA11::1 301
frame-r map ipv6 FE80::2 302
frame-r map ipv6 FE80::1 301
no frame-relay inverse-arp
no shut
```

14. IPv6 Access Control Lists

Lista de control de acceso (ACL), es una lista de reglas que detallan puertos de servicio o nombres de dominios (de redes) que están disponibles en un terminal u otro dispositivo de capa de red, cada uno de ellos con una lista de terminales y/o redes que tienen permiso para usar el servicio.

Tantos servidores individuales como enrutadores pueden tener ACL. Las listas de control de acceso pueden configurarse generalmente para controlar tráfico entrante y saliente y en este contexto son similares a un cortafuego (firewall).

14.1 ACLs de entrada y salida

Las ACLs actúan únicamente sobre el tráfico que proviene de un enlace conectado a una interfaz, nunca sobre el tráfico que se ha generado dentro del propio router.

Para filtrar eficazmente el tráfico, se configuran de dos maneras:

ACLs de entrada: Los paquetes se filtran cuando intentan entrar en el router a través de una interfaz, antes de ser procesados por el router. Esto evita que se procese tráfico innecesario y se gane rendimiento en los routers.

ACLs de salida: Una vez que ha entrado el paquete en el router y ha sido procesado y elegida la interfaz de salida según su tabla de enrutamiento, se llevará a cabo el filtrado de la ACL.

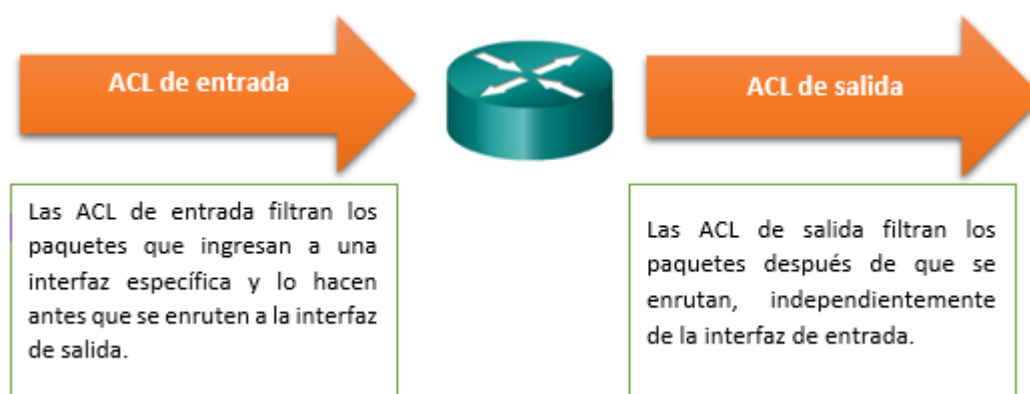


Figura 91: ACL de entrada y salida.

Las listas de acceso determinan qué tráfico está bloqueado y qué tráfico se reenvía a las interfaces del dispositivo y permiten filtrado de tráfico basado en direcciones de origen y de destino, y el tráfico entrante y saliente a una interfaz específica. La funcionalidad estándar IPv6 ACL se amplió para apoyar filtrado de



tráfico basado en la opción de cabeceras de IPv6 y, de capa superior de información de tipo de protocolo para mejorar gradualmente su control. Las listas de acceso tienen su enfoque en:

- Limitar el tráfico de red y mejorar el rendimiento de la red.
- Brindar control de flujo de tráfico.
- Proporcionar un nivel básico de seguridad para el acceso a la red.
- Se debe decidir qué tipos de tráfico se envían o bloquean en las interfaces del router.

Existen varias diferencias entre las ACL de IPv6 y las de IPv4:

- Al aplicar la lista en una interfaz: en IPv4 se usa el comando **ip access-group** pero en IPv6 se usa **ipv6 traffic-filter**.
- No hay máscaras wildcard, sino que se usa el prefijo de la dirección IPv6 para indicar cuánto de la dirección debe coincidir.
- Hay dos sentencias implícitas más aparte del **deny any any** de las listas extendidas de IPv4, hay dos sentencias más que se aplican automáticamente si no hay ninguna coincidencia previa:

- **permit icmp any any nd-na** – Neighbor Advertisement.

- **permit icmp any any nd-ns** – Neighbor Solicitation.

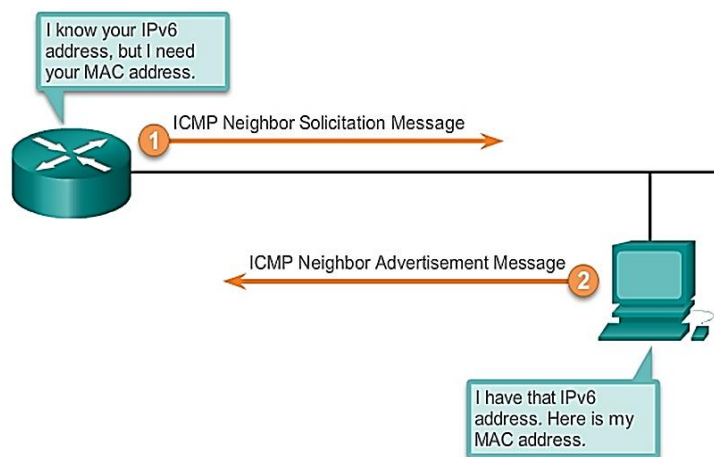


Figura 92: Sentencias para consultar mac en IPv6

Estas sentencias sirven para que el router pueda consultar las MAC de los vecinos a los que les va a mandar los paquetes (similar a ARP en IPv4). Este proceso se hace en capa 3 en IPv6 mientras que en IPv4 se realiza en la capa 2.



14.2 Paquetes de inspección en IPv6

Los siguientes campos de cabecera se utilizan para la inspección IPv6: clase de tráfico, etiqueta de flujo, longitud de carga útil, la siguiente cabecera, límite de saltos, y de origen o destino de dirección IP.

14.3 Limitaciones ACL en IPv6

El proceso de creación de ACLs en IPv6 es algo más sencillo que en IPv4, porque ni hay máscaras wildcard ni hay tantos tipos. De hecho, solo hay un tipo de ACL, equivalente a la extendida de nombre en IPv4.

Las ACL extendidas filtran el tráfico en función del tipo de protocolo que usen, sus direcciones IP de origen y destino y los puertos TCP o UDP de origen y destino.

Es importante tener en cuenta que las ACLs de IPv4 e IPv6 no pueden tener el mismo nombre.

Las direcciones IPv6 origen y destino ACL sólo son compatibles con los prefijos de /0 a /64 y las direcciones de host (/128) que se encuentran en el identificador universal, extendido (EUI) Formato de 64. El switch soporta solamente estas direcciones de host sin pérdida de información:

- Direcciones unicast globales agregables.
- Direcciones Link-local

El switch no soporta estas palabras clave: control de etiquetas, cabecera de enrutamiento, e indeterminado-transporte. Esta versión sólo admite ACL y ACL puerto del router de IPv6; que no soporta VLAN ACL (mapas de VLAN). El switch no aplica ACL basadas en MAC en las tramas de IPv6. No se pueden aplicar las ACL en puerto IPv6 de Capa 2 EtherChannels. El switch no soporta ACL con puertos de salida. ACL del router de salida y de entrada de ACL del puerto para IPv6 sólo se admiten en las pilas del switch. Los switch sólo admiten control de flujo (entrante) IPv6 ACL. Al configurar una ACL, no hay ninguna restricción en palabras claves introducidas en la ACL, independientemente de si son o no son compatibles con la plataforma. Al aplicar la ACL a una interfaz que requiere el reenvío de hardware (puertos físicos o SVI), los switches hacen un chequeo para determinar si la ACL es soportada por la interfaz, si no son directamente rechazadas.

Si una ACL se aplica a una interfaz y se intenta añadir una entrada de control de acceso (ACE) con una palabra clave no compatible, cada comando se denomina ACE. Uno o más ACE con el mismo nombre de lista se conocen como una lista de acceso, el switch no permite ACE para ser añadida a la ACL que se une actualmente a la interfaz.



14.4 Configuración de ACL en IPv6.

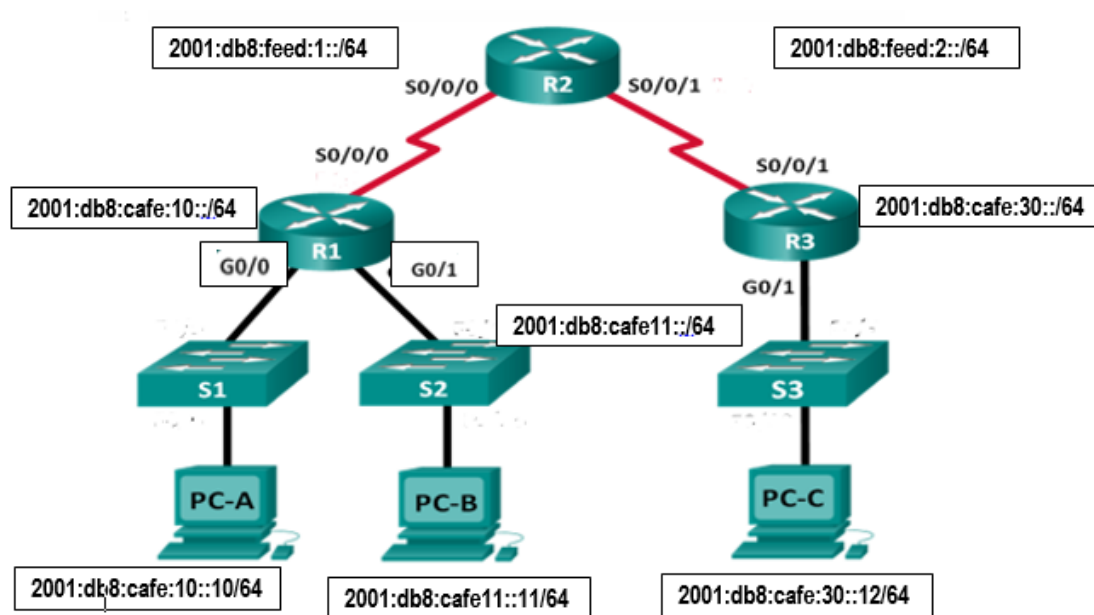


Figura 93: Topología de ACL.

Crear la ACL

Vamos a crear una ACL llamada ACCESO-RESTRINGIDO que bloquee el tráfico procedente de la red 2001:db8:cafe:30::/64 pero permita todo lo demás:

```
R1(config)# ipv6 access-list ACCESO-RESTRINGIDO
R1(config-ipv6-acl)# deny ipv6 2001:db8:cafe:30::/64 any
R1(config-ipv6-acl)# permit ipv6 any any
```

También podemos capar el acceso por telnet de cualquier sitio hacia el host 2001:db8:cafe:11::11 con:

```
R1(config-ipv6-acl)# deny tcp any host 2001:db8:cafe:11::11 eq 23
```



Aplicar la ACL a una interfaz

Vamos a aplicar la ACL que acabamos de crear en la interfaz serial 0/0/0 como entrada:

```
R1(config)# interface s0/0/0
R1(config-if)# ipv6 traffic-filter ACCESO-RESTRINGIDO in
```

Verificar ACLs en IPv6

ACLs aplicadas una interfaz concreta.

Con el comando `show ipv6 interface g0/0` podemos ver, por ejemplo, las listas aplicadas a la interfaz gigabitEthernet 0/0.

```
R3# show ipv6 interface g0/0
GigabitEthernet0/0 is up, line protocol is up
Global unicast address(es):
  2001:DB8:CAFE:30::1, subnet is 2001:DB8:CAFE:30::/64
Input features: Access List
Inbound access list RESTRICTED-ACCESS
```

Listas creadas en el router

Podemos mostrar las listas creadas en el router y las sentencias que las componen con el comando `show access-lists`.

```
R3# show access-lists
IPv6 access list RESTRICTED-ACCESS
 permit tcp any host 2001:DB8:CAFE:10::10 eq www sequence 20
 permit tcp any host 2001:DB8:CAFE:10::10 eq 443 sequence 30
 deny ipv6 any 2001:DB8:CAFE:10::/64 sequence 50
 permit tcp host 2001:DB8:CAFE:30::12 host 2001:DB8:CAFE:11::11 eq
 telnet sequence 70
 deny tcp any host 2001:DB8:CAFE:11::11 eq telnet sequence 90
 permit ipv6 any any sequence 110
```



15. Mecanismos de transición.

La conversión de redes IPv4 a IPv6 tardará un largo período de tiempo, por lo que en el diseño de IPv6 se han tomado en cuenta mecanismos que permitan la coexistencia y comunicación de ambos protocolos. Se han diseñado mecanismos (RFC 1933) que permiten la coexistencia de ambos protocolos.

Destacan:

- Técnicas de Dual Stack (Doble Pila). Permiten a IPv4 y a IPv6 coexistir en los mismos dispositivos y redes.
- Técnicas de Tunneling. Permiten el transporte de tráfico de IPv6 a través de la infraestructura de IPv4 existente.
- Técnicas de traducción. Permiten comunicar solamente nodos IPv6 con nodos IPv4.

15.1 DSTM

El DSTM (Dual Stack Transition Mechanism) se compone de dos métodos en particular: AIH (Assignment of IPv4 global addresses to IPv6 hosts) y DTI (Dynamic Tunneling Interface). AIH es un método que permite asignar temporalmente direcciones IPv4 a hosts Dual Stack dentro de una red IPv6. DTI es una interface diseñada para encapsular paquetes IPv4 dentro de paquetes IPv6. La unión de ambos métodos da como resultado el mecanismo DSTM, el cual tiene como objetivo que un host IPv6 obtenga una dirección IPv4 para establecer comunicación con hosts que manejen exclusivamente direcciones IPv4.

DSTM permite también ejecutar aplicaciones IPv4 sin modificación alguna, y solo se puede aplicar dentro de una red IPv6. A continuación se muestra un esquema del mecanismo DSTM:

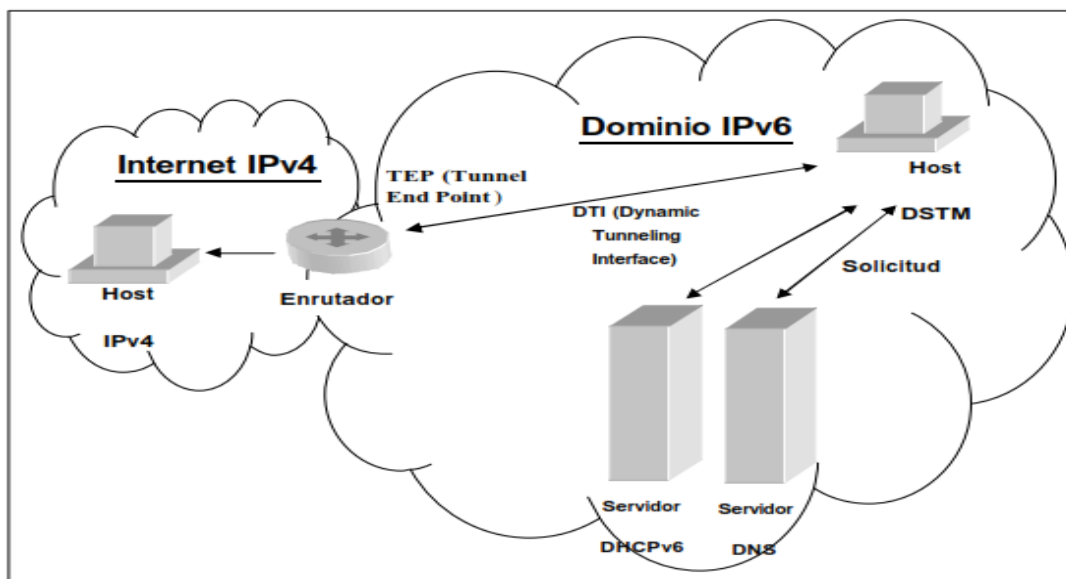


Figura 94: Entorno DSTM

El entorno DSTM trabaja solamente con hosts Dual Stack. Se necesita un servidor encargado de asignar temporalmente direcciones IPv4 a los hosts, el cual generalmente utiliza DHCPv6, ya que DHCPv4 no puede ser utilizado dentro de una red IPv6. También se necesita un servidor para resolución de DNS y un enrutador frontera con soporte Dual Stack para comunicar el dominio IPv6 a un dominio exterior o al Internet.

Funcionamiento de DSTM.

El host DSTM hace una solicitud de dirección IPv4 al servidor de direcciones, para establecer comunicación con un host fuera del dominio. Para establecer comunicación con un host dentro del mismo dominio, no es necesario solicitar una dirección IPv4.

El servidor de direcciones le asigna una dirección IPv4 temporalmente al host DSTM. El tiempo de vida de esa asignación debe ser indicado en la respuesta del servidor al host, así como la dirección IPv6 del TEP (Tunnel End Point). Si un host requiere más tiempo, la dirección IPv4 tendrá que completar el tiempo y realizar una nueva solicitud al servidor de direcciones. El servidor de direcciones también se encargará de mapear las direcciones IPv4 asignadas a la dirección IPv6 correspondiente, es decir, relacionar ambas direcciones. Como éstas direcciones son asignadas temporalmente, se podrán guardar en una memoria cache. Como una extensión del proceso de asignación de direcciones, el servidor puede asignar un rango de puertos a utilizar por el host. Esto permite que una sola dirección IPv4 pueda ser utilizada por varios hosts al mismo tiempo, evitando que los puertos se traslapen.

Con la dirección IPv6 del TEP proporcionada por el servidor de direcciones, el host se encargará de configurar una interfase DTI hacia el TEP, encapsulando los paquetes IPv4 dentro de paquetes IPv6. Si la



interfase no ha sido configurada, es decir, que no tiene asignada una dirección IPv4, el proceso deberá detenerse hasta obtener una dirección IPv4 del servidor de direcciones. Todo el tráfico IPv4 puede ser dirigido a esta interfase por medio de una entrada en la tabla de enrutamiento del host. Una vez que la dirección IPv4 ha sido asignada, es utilizada como dirección fuente para todos los paquetes que sean enviados desde esa interfase.

Por último, el host manda los paquetes encapsulados hacia el TEP, generalmente el enrutador frontera. Este se encarga de decapsular los paquetes y reenviarlos hacia la red exterior o el Internet, de modo que lleguen al host solicitado por el host Dual Stack.

❖ **Comunicación Bidireccional**

El mecanismo DSTM es bidireccional, es decir, permite que un host Dual Stack dentro de un dominio IPv6 se comunique con hosts exclusivamente IPv4, o en caso contrario, que un host exclusivamente IPv4 se pueda comunicar con un host Dual Stack dentro de un dominio IPv6.

En el primer caso, el host Dual Stack solicitará una resolución de dirección de tipo AAAA para el host con el que quiere establecer comunicación. Debido a que el host no es IPv6, el servidor DNS le devolverá un error de resolución. Es entonces cuando el host Dual Stack solicitará una dirección IPv4 para poder establecer la comunicación.

En el segundo caso, cuando un host exclusivamente IPv4 desea establecer comunicación con un host Dual Stack dentro de un dominio IPv6, el host IPv4 solicita una resolución de dirección de tipo A para el host Dual Stack al servidor DNS dentro de su red IPv4. El servidor DNS se comunica con el enrutador DS TM, el cual solicita al servidor de direcciones del dominio IPv6 que le asigne temporalmente una dirección IPv4 al host solicitado, para poder establecer la comunicación.

15.2 Túneles

❖ **6to4**

Este método es también conocido como “Connection of IPv6 Domains vía IPv4 Clouds” (conexión de dominios IPv6 por medio de nubes IPv4). Esencialmente, este método permite a sitios o hosts IPv6 comunicarse entre ellos a través de una red IPv4, sin necesidad de configuración manual de túneles, y permite que dichos sitios o hosts se comuniquen con el Internet IPv6 por medio de enrutadores 6to4 Relay.

Este método debe ser temporal, y se utilizará mientras se obtenga una conexión IPv6 nativa, es decir, mientras se lleva a cabo la transición de IPv4 a IPv6. No fue diseñado como una solución permanente.

El esquema de este método se muestra en la figura siguiente:

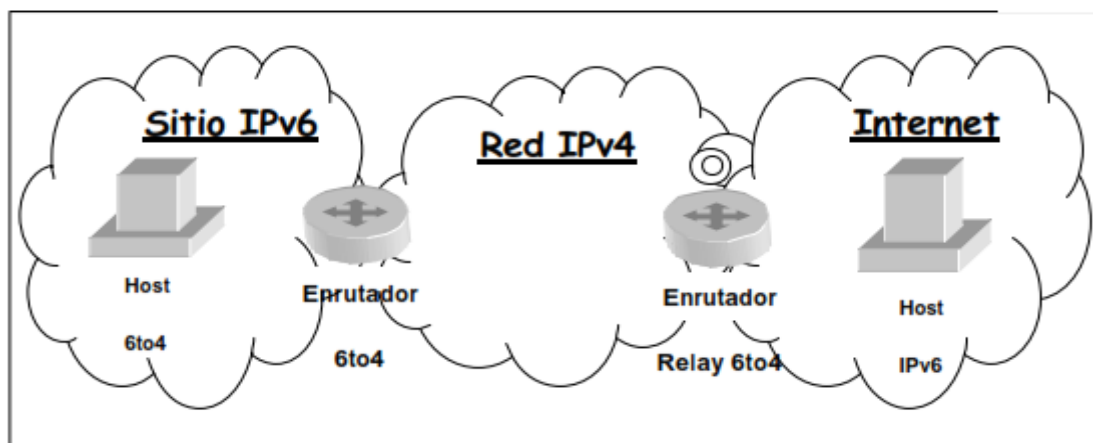


Figura 95: Entorno 6to4

Dentro de este entorno encontramos principalmente 3 elementos:

- Un host 6to4 es un host IPv6 que tiene configurada al menos una dirección de tipo 6to4. Estos hosts no requieren configuración manual, comúnmente cuentan con un mecanismo de autoconfiguración.
- Un enrutador 6to4 es un enrutador Dual Stack que soporta el uso de túneles 6to4, el cual sirve para intercambiar paquetes de tipo 6to4 entre enrutadores del mismo tipo y sitios o hosts IPv6. Estos enrutadores requieren configuración manual adicional, ya que son los encargados de encapsular y decapsular los paquetes.
- Un enrutador 6to4 Relay se puede definir como un enrutador 6to4 configurado para soportar el enrutamiento de tránsito entre direcciones 6to4 y direcciones IPv6 nativas. Este enrutador debe tener al menos una interfase 6to4 y una interfase IPv6 nativa, para poder establecer comunicación entre dominios IPv4 e IPv6.

Dirección 6to4

Una dirección de tipo 6to4, está conformada por distintas partes como se muestra en la siguiente figura:

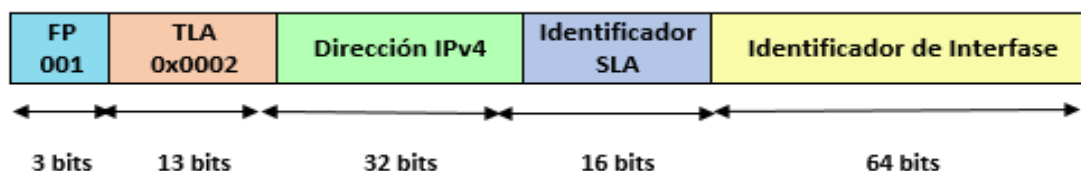


Figura 96: Dirección 6to4



Este tipo de dirección utiliza el prefijo 001, el cual identifica a las direcciones Unicast Globales Agregables, seguido por un identificador TLA de 13 bits asignado por IANA, cuyo valor es 0x0002. Después le sigue la dirección IPv4 del sitio, así como un identificador SLA y el identificador de interfase.

Todo esto se puede expresar 2002:DirecciónIPv4::/48. La forma en que se convierte la dirección IPv4 al formato para estas direcciones se muestra a continuación:

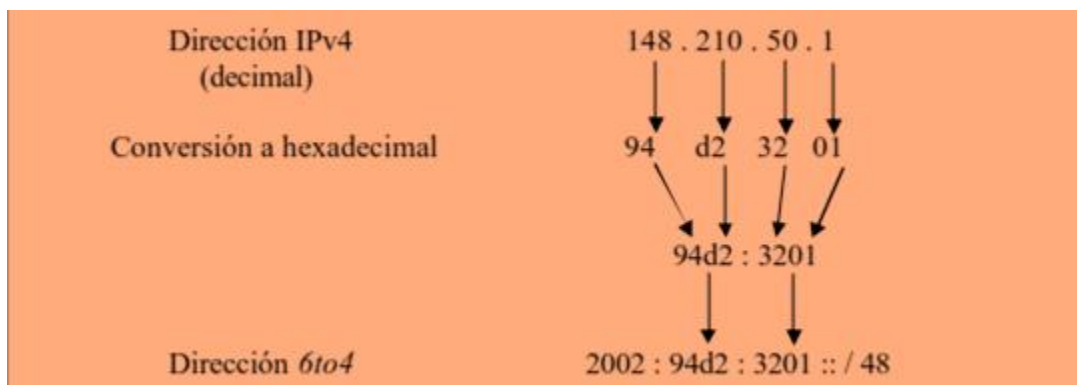


Figura 97: Conversión de dirección IPv4 a dirección 6to4

Selección de dirección

En caso de que un host tenga una dirección 6to4, y el host con el que quiere establecer comunicación tenga una dirección 6to4 y una dirección IPv6 nativa, es recomendable que ambos hosts establezcan la comunicación utilizando 6to4. En caso de que ambos hosts tengan direcciones 6to4 y direcciones IPv6 nativas, se puede establecer la comunicación siempre y cuando ambos hosts utilicen el mismo tipo de direcciones, aunque es recomendable que la comunicación se realice por medio de direcciones IPv6 nativas.

Encapsulación 6to4

En el método de 6to4 se utiliza la encapsulación de paquetes IPv6 dentro de paquetes IPv4. El campo "Protocolo" de la cabecera IPv4 debe ser igual a 41, que es el número asignado para este tipo de encapsulación o túneles. Las direcciones de destino y origen, ubicadas en la cabecera IPv4, pueden ser las mismas direcciones del campo que contiene la dirección IPv4 en el prefijo formado para las direcciones 6to4.

Tipos de comunicación

Los enrutadores IPv6 dentro de un mismo sitio publican prefijos 2002:direcciónIPv4:identificadorSLA::/64 para permitirle a los hosts crear direcciones 6to4 autoconfiguradas. Los hosts o subredes individuales se configuran automáticamente con una ruta de 64 bits de una subred para intercambio directo entre hosts

vecinos. Cualquier paquete IPv6 que no contenga un prefijo de 64 bits similar al de alguna de las subredes del sitio, será enviado al enrutador 6to4 colocado en la frontera del sitio.

Con el método de 6to4 se pueden efectuar varios tipos de comunicación. A continuación se muestra un entorno para ejemplificar los diferentes tipos de comunicación:

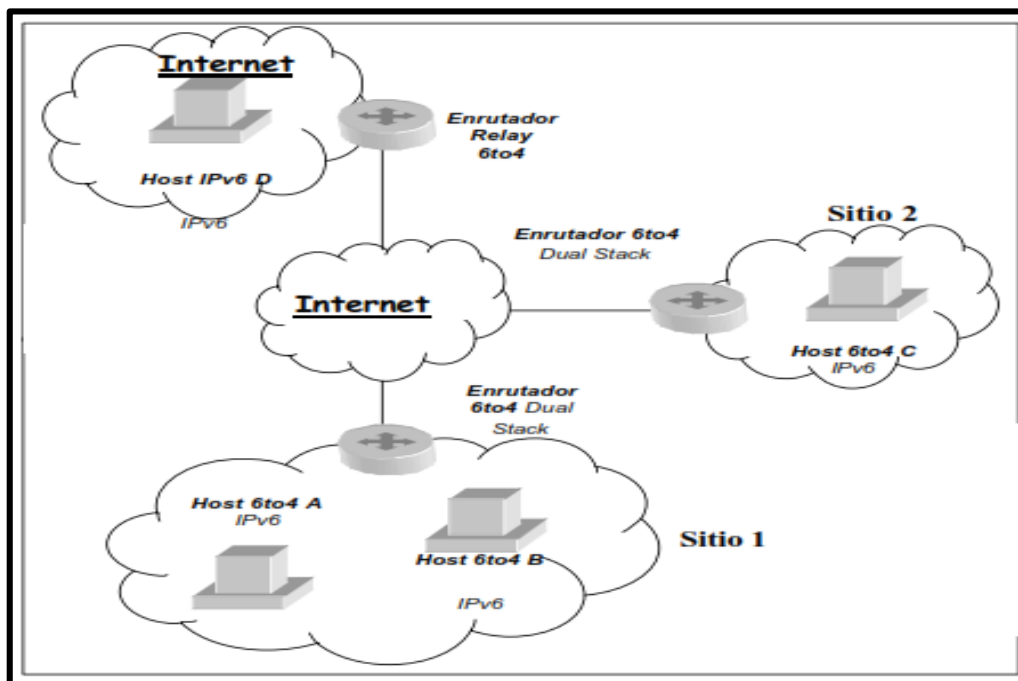


Figura 98: Comunicación 6to4

- **Comunicación del Host A al Host B.** Un host 6to4 puede establecer comunicación con un host 6to4 dentro de su mismo sitio. El host origen (Host A) envía el paquete al host solicitado (Host B) utilizando la infraestructura del sitio local (Sitio 1).
- **Comunicación del Host A al Host C.** Un host 6to4 puede establecer comunicación con hosts 6to4 en otros sitios. Lo primero que hace el host origen (Host A) es mandar el paquete al enrutador 6to4 local (Sitio 1). Después este se encarga de hacerlo llegar al enrutador 6to4 del sitio solicitado (Sitio 2) por medio de la creación de túneles en la infraestructura IPv4. Por último, el enrutador en el sitio destino (Sitio 2) se encarga de decapsular el paquete y entregarlo al host solicitado (Host C) utilizando la infraestructura IPv6 del sitio.
- **Comunicación del Host A al Host D.** Un host 6to4 puede establecer comunicación con hosts en el Internet IPv6. Lo primero que hace el host origen (Host A) es mandar el paquete al enrutador 6to4 local (Sitio 1). Después este se encarga de hacerlo llegar a un enrutador 6to4 Relay, el cual tenga acceso a ambos entornos, Internet IPv4 e Internet IPv6. Por último, el enrutador 6to4 Relay se encarga de

desencapsular el paquete y entregarlo al host solicitado (Host D) utilizando la infraestructura IPv6 del sitio

❖ 6over4

Este método permite que hosts IPv6 que se encuentren dentro de un dominio IPv4, y que no están conectados directamente a un enrutador Dual Stack, establezcan una comunicación con otros hosts IPv6 dentro del mismo dominio mediante la encapsulación de paquetes IPv6 dentro de paquetes IPv4. Si alguno de estos hosts IPv6 desea establecer comunicación con algún host ubicado en otro dominio IPv6, será necesario que exista de por medio un enrutador Dual Stack.

A este método se le conoce formalmente como “IPv6 over IPv4”, pero comúnmente se le conoce como “6over4” o “Virtual Ethernet” (se entiende como una capa de enlace virtual). El esquema de este método se muestra en la figura siguiente:

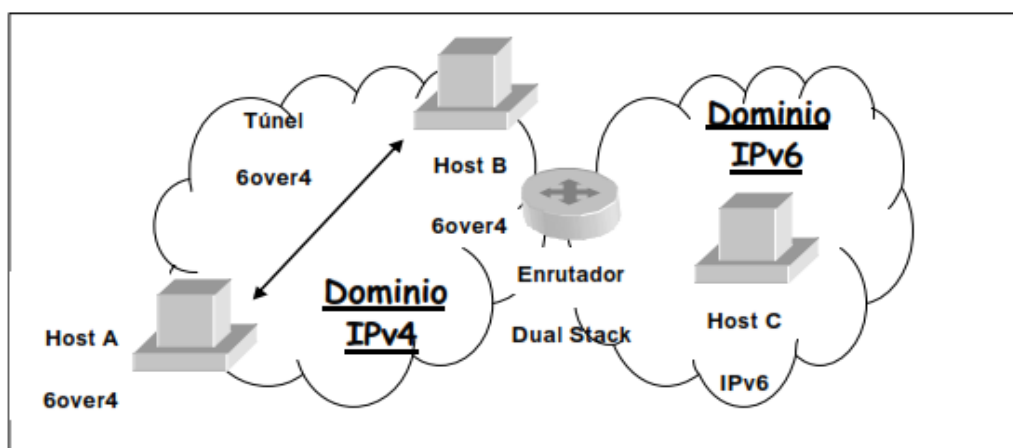


Figura 99: Entorno 6over4

Dirección 6over4

El dominio IPv4 debe ser multicast para que se puedan llevar a cabo algunos mensajes o procesos de descubrimiento de nodos vecinos (**Neighbor Discovery**). Para la traducción de direcciones IPv6 multicast a direcciones IPv4 multicast se ha definido el siguiente patrón:

239.192.[segundo byte más a la derecha de la dir. IPv6].[último byte de la dir. IPv6]

A continuación se muestran algunos ejemplos de direcciones IPv6 multicast traducidas a direcciones IPv4 multicast:

- FF02::1 (dirección multicast de enlace local con alcance a todos los hosts) se cambia por la dirección 239.192.0.1



- FF02::2 (dirección multicast de enlace local con alcance a todos los enrutadores) se cambia por la dirección 239.192.0.2
- FF02::1:FF45:8C54 (dirección multicast de un nodo solicitado) se cambia por la dirección 239.192.140.84

Cuando se utiliza 6over4, el dominio IPv4 hace uso de mensajes **IGMP** (Internet Group Management Protocol) para informar a los enrutadores IPv4 locales del tráfico multicast que está siendo enviado. Los hosts que soportan 6over4 también registran direcciones MAC multicast adicionales para sus adaptadores de red, y estas son correspondientes a las direcciones IPv4 multicast. A continuación se muestran algunas direcciones MAC

Correspondientes a direcciones IPv4 multicast para un adaptador de tipo Ethernet:

La dirección MAC multicast correspondiente a la dirección 239.192.0.1 es 01-00-5E-40-00-01.

La dirección MAC multicast correspondiente a la dirección 239.192.0.2 es 01-00-5E-40-00-02.

La dirección MAC multicast correspondiente a la dirección 239.192.140.84 es 01-00-5E-40-8C-54

MTU

La unidad máxima de transmisión o MTU (Maximum Transmission Unit) por default de los paquetes IPv6 en un dominio IPv4 deberá ser de 1480 octetos (la MTU máxima normal es de 1500 octetos, pero se deben reservar 20 octetos para la cabecera IPv4). Este tamaño puede variar cuando se especifique alguna MTU en un Router Advertisement o por alguna configuración manual. En el caso de que la MTU sea muy grande para una red intermedia, esto asegurará una fragmentación, por lo que en este caso se debe asegurar que el bit de "DF" de la cabecera IPv4 no este activo.

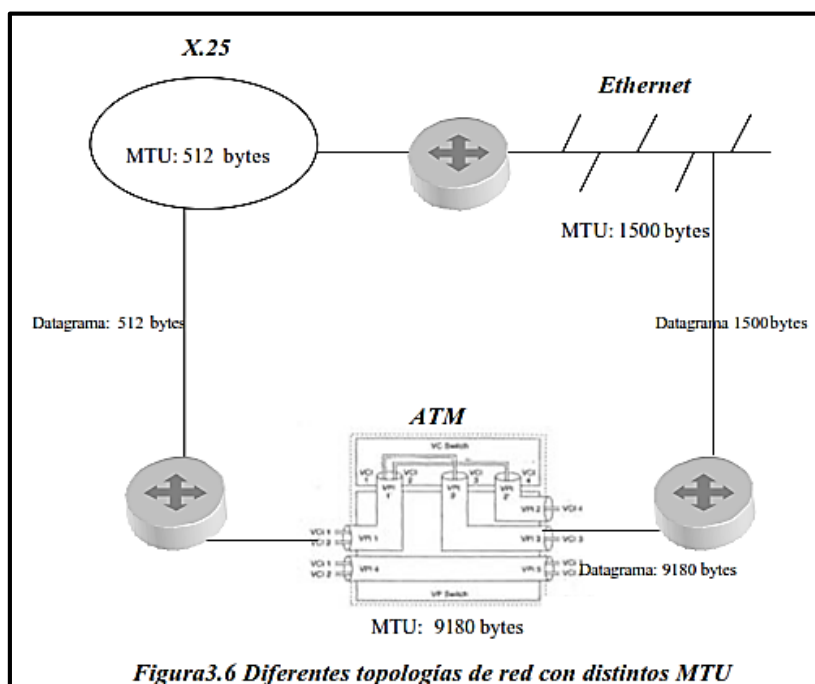


Figura 100: Diferentes topologías de red con distintos MTU.

Encapsulación 6over4

Los paquetes IPv6 serán transmitidos dentro de paquetes IPv4 con el campo de "Protocolo" igual a 41, que es el valor predefinido para túneles o encapsulación de paquetes IPv6 en paquetes IPv4. La cabecera IPv4 contiene las direcciones IPv4 de origen y destino. El cuerpo del paquete IPv4 contiene la cabecera IPv6, así como su carga. En caso de que el paquete IPv4 contenga opciones, estas deberán ser rellenas hasta terminar todo el bloque de 32 bits, y que la cabecera IPv6 comience en un nuevo bloque de 32 bits.

❖ Tunnel Broker

Desde el comienzo de IPv6 y durante su crecimiento, se ha tenido la necesidad de utilizar la infraestructura de red existente, es decir, la infraestructura IPv4. La mayoría del 6bone está conectado por una variedad de túneles de distintos tipos, cada uno de ellos con un objetivo en especial, pero al mismo tiempo con ciertos problemas o limitaciones.

Tipos de túneles

Los siguientes son los principales tipos de túneles:

- **Túneles automáticos con direcciones IPv4 compatibles.**

Útiles para conectar enrutadores o hosts que se encuentran aislados, pero trae consigo el problema de la escasez de direcciones IPv4 que prácticamente es el problema a solucionar. Además, las tablas de enrutamiento IPv4 seguirán creciendo cada vez más, y lo peor de todo es que estas direcciones se tendrán que almacenar también en las tablas de enrutamiento IPv6, lo cual creará un grande problema con respecto al tamaño de dichas tablas.



- **Túneles de tipo 6to4.**

Permiten a dominios IPv6 aislados que cuenten con una conexión directa a una red IPv4 o al Internet, poder establecer comunicación con otros dominios IPv6 con una mínima configuración manual. Este tipo de túneles se utiliza comúnmente en redes aisladas o privadas.

- **Túneles de tipo 6over4.**

Es un mecanismo a nivel de sitio que utiliza un dominio IPv4 Multicast, como una capa de enlace de datos virtual. Sin embargo necesita un enrutador extra si desea establecer comunicación con un dominio IPv6 externo.

Los túneles manualmente configurados han sido de gran ayuda hasta la fecha, pero requieren una estricta supervisión y mantenimiento de parte de los administradores de las redes, por lo que se pensó en crear un método que creara túneles configurados de una manera automática. Es aquí donde nació el concepto de Tunnel Broker (TB).

Descripción de TB

La idea principal de este método es tener servidores dedicados, llamados TB's, que se encarguen de configurar túneles de una manera automática en respuesta a requisiciones hechas por los usuarios de este servicio, y de esta manera aumentar el número de hosts que se encuentran actualmente conectados a una red IPv6. Se espera que en un futuro existan varios tipos de TB's, de manera que el usuario pueda seleccionar de una lista el que mejor se acomode a sus necesidades, por ejemplo el más cercano, el más barato, etc. El método de TB permite a hosts Dual Stack que cuenten con una conexión a una infraestructura IPv4, crear túneles automáticamente para poder establecer una comunicación con dominios IPv6.

A continuación se muestra el entorno de un TB de una manera gráfica:

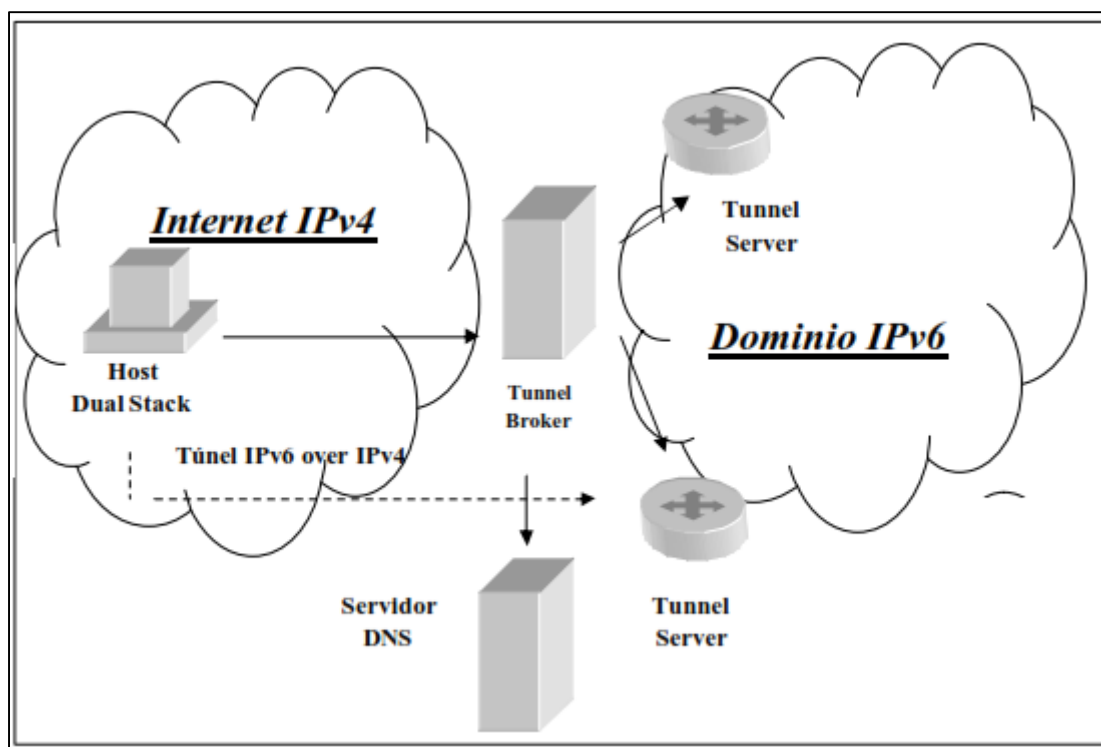


Figura 101: Entorno Tunnel Broker

Encapsulación 6over4

Los paquetes IPv6 serán transmitidos dentro de paquetes IPv4 con el campo de "Protocolo" igual a 41, que es el valor predefinido para túneles o encapsulación de paquetes IPv6 en paquetes IPv4. La cabecera IPv4 contiene las direcciones IPv4 de origen y destino. El cuerpo del paquete IPv4 contiene la cabecera IPv6, así como su carga. En caso de que el paquete IPv4 contenga opciones, estas deberán ser rellenadas hasta terminar todo el bloque de 32 bits, y que la cabecera IPv6 comience en un nuevo bloque de 32 bits.

15.3 Traductores

➤ SIIT

El método de SIIT (Stateless IP/ICMP Translation algorithm) básicamente se encarga de traducir las cabeceras entre IPv4 e IPv6 (incluyendo las cabeceras ICMP), y permite la comunicación entre hosts exclusivamente IPv6 y hosts exclusivamente IPv4. El nodo IPv6 de alguna forma obtendrá una dirección IPv4 temporal y un medio de enrutamiento para los paquetes. La dirección IPv4 temporal será utilizada como una dirección IPv6 llamada IPv4 - traducida. Después los paquetes pasarán por un traductor SIIT encargado de traducir las cabeceras de los paquetes IPv4 e IPv6, así como las direcciones en las

cabeceras. Las direcciones utilizadas en este método pueden ser IPv4, IPv4-traducidas o IPv4-mapeadas. Este método no especifica de qué manera se obtendrá la dirección IPv4 temporal (se sugiere DHCP con algunas extensiones), ni cómo será registrada en el DNS. Tampoco especifica el tipo de enrutamiento de los paquetes.

El método de SIIT puede ser utilizado cuando se desea establecer comunicación entre redes IPv6 pequeñas o hosts IPv6 y hosts IPv4, como se muestra en la siguiente figura:

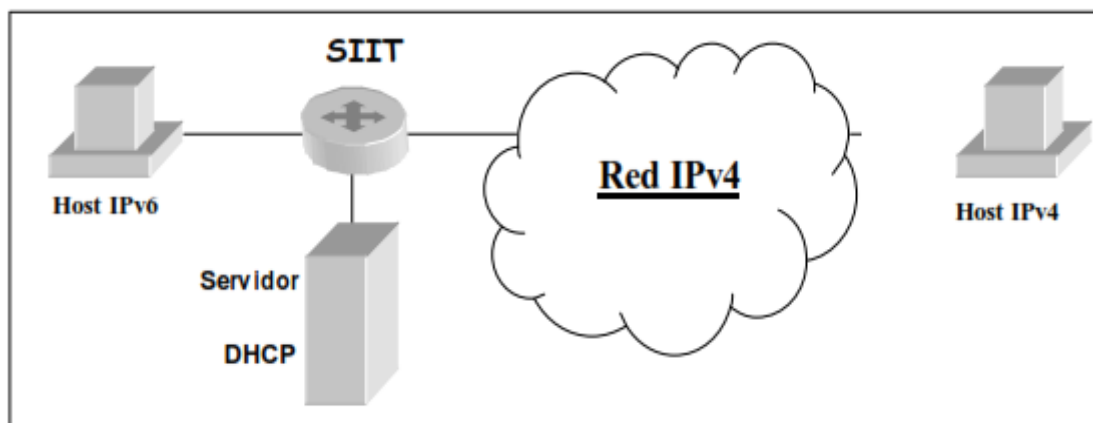


Figura 102: SIIT para redes pequeñas IPv6.

También se puede utilizar cuando se desea establecer comunicación entre hosts IPv6 en una red Dual Stack y hosts IPv4, como se muestra en la siguiente figura:

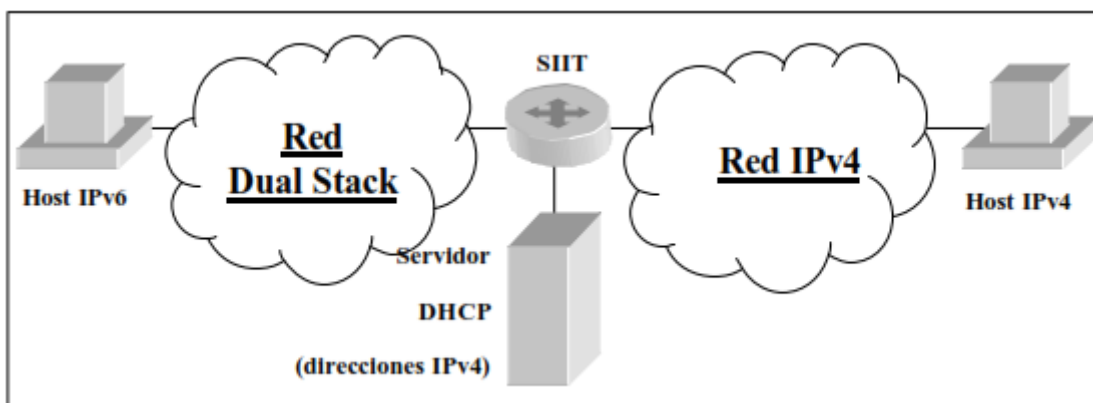


Figura 103: SIIT para redes Dual Stack.

Este método no es recomendable después de la transición, ya que solo existirán algunas redes IPv4 pequeñas y los traductores se encontrarían en los límites de estas, lo que significa un largo recorrido de los paquetes provenientes de los *hosts* IPv6 para obtener una dirección IPv4 temporal, la cual les permitiría llevar a cabo la comunicación.



Las direcciones utilizadas por este método son las siguientes:

- IPv4- mapeada.- Una dirección de la forma 0::FFFF:a.b.c.d que identifica a un nodo que no soporta IPv6.
- IPv4-traducida.- Una dirección de la forma 0::FFFF:0:a.b.c.d que identifica a un nodo que soporta IPv6.

✓ Traducción de IPv4 a IPv6

Cuando un traductor IPv4-IPv6 recibe un paquete IPv4 destinado a un host que se encuentra fuera de su dominio, debe traducir la cabecera IPv4 del paquete a una cabecera IPv6, para después reenviar ese paquete al exterior del dominio. La cabecera IPv4 es removida completamente y reemplazada por la cabecera IPv6. Los paquetes ICMP, así como la cabecera de transporte y la carga de datos no cambian. En IPv6 es necesario realizar un Path MTU Discovery antes de enviar un paquete, pero en IPv4 no lo es. Esto significa que los enrutadores intermedios IPv6 nunca fragmentan paquetes, el único habilitado para hacer esto es el host que manda el paquete.

Cuando un host IPv4 realiza un Path MTU Discovery (habilitando el bit “DF” de la cabecera), puede realizar este proceso de host a host pasando por un traductor. En este proceso, enrutadores IPv4 o IPv6 podrán enviar mensajes ICMP de vuelta al host para indicarle que los paquetes son muy grandes (“packet too big”). Cuando estos paquetes sean enviados por un enrutador IPv6, estos deberán pasar por el traductor para que este se encargue de efectuar la traducción de paquetes ICMP a una forma en que el host IPv4 pueda entenderlos.

En caso contrario, cuando un host IPv4 no realiza un Path MTU Discovery, el traductor debe asegurar que los paquetes no rebasan el MTU máximo del enlace IPv6. Esto se logra mediante la fragmentación de los paquetes IPv4 de una manera que quepan en paquetes IPv6 de 1280 bytes, ya que con este tamaño se garantiza que los paquetes no tendrán que ser fragmentados.

✓ Traducción de cabeceras IPv4 a cabeceras IPv6

Los campos de la cabecera IPv6 se traducen como sigue:

Versión	Version	6
Clase de Tráfico	Traffic Class	Se copia del campo Type of Service de la cabecera IPv4.
Etiqueta de flujo	Flow Label	0 (todos los bits en cero).



Longitud de carga útil	Payload Length	Longitud total de la cabecera IPv4, menos el tamaño de esta y sus opciones, si existen.
Próxima Cabecera	Next Header	Se copia del campo Protocol de la cabecera IPv4.
Límite de saltos	Hop Limit	Se copia del campo Time to Live de la cabecera IPv4. Como el traductor es un enrutador, este debe decrementar el campo Time to Live de la cabecera IPv4 o Hop Limit de la cabecera IPv6 antes de reenviar el paquete. Después de decrementar el valor también debe verificar que este no esté en cero. Si se encuentra en cero, deberá mandar un mensaje de error ("ttl exceeded").
Dirección de origen	Source Address	Los 32 bits más a la derecha son la dirección fuente IPv4. Los 96 bits anteriores son el prefijo para las direcciones IPv4- mapeadas (::FFFF:0:0/96).
Dirección de destino	Destination Address	Los 32 bits más a la derecha son la dirección destino IPv4. Los 96 bits anteriores son el prefijo para las direcciones IPv4- traducidas (0::FFFF:0:0/96).

Tabla 22: Traducción de cabeceras IPv4 a cabeceras IPv6

Si se encuentran opciones en la cabecera IPv4, estas deberán ser ignoradas, no deberán ser traducidas. Si fuera necesario agregar una cabecera de fragmentación (el bit de "DF" no está activo o el paquete es un fragmento), los campos se deben traducir como sigue:

Campos de la cabecera IPv6:

Longitud de carga útil	Payload Length	Longitud total de la cabecera IPv4, mas 8 de la cabecera de fragmentación, menos el tamaño de la cabecera IPv4 y sus opciones, si existen.
Próxima Cabecera	Next Header	Cabecera de fragmentación (44).

Tabla 23: Campos de la cabecera IPv6



Campos de la cabecera de fragmentación IPv6:

Próxima Cabecera	Next Header	Se copia del campo Protocol de la cabecera IPv4.
Compensación de Fragmento de cabecera	Fragment Offset	Se copia del campo Fragment Offset de la cabecera IPv4.
Bandera M	M Flag	Se copia el bit del campo More Fragments de la cabecera IPv4.
Identificación	Identification	Se copian los 16 bits más a la derecha del campo Identification de la cabecera IPv4. Los 16 bits restantes se rellenan con ceros.

Tabla 24: Campos de la cabecera de fragmentación IPv6

✓ Traducción de IPv6 a IPv4

Cuando un traductor IPv6-IPv4 recibe un paquete IPv6 destinado a una dirección IPv6 de tipo IPv4-mapeada, debe traducir la cabecera IPv6 a una cabecera IPv4, para después reenviar ese paquete a su destino. La cabecera IPv6 original es removida completamente y reemplazada por la cabecera IPv4. Los paquetes ICMP, así como la cabecera de transporte y la carga de datos no cambian.

Un enlace IPv6 debe tener un MTU de 1280 bytes o más, mientras que IPv4 debe tener un MTU de 68 bytes. Entre IPv4 e IPv6 existen diferencias que afectan la traducción, tales como la fragmentación y el MTU de los enlaces. No es posible realizar un Path MTU Discovery de host a host cuando existe un traductor IPv6-IPv4, debido a que el host IPv6 puede recibir mensajes ICMP de error, indicándole que los paquetes son muy grandes originados por un enrutador IPv4 que reporte un MTU menor de 1280 bytes.

Los host IPv6 responden a estos mensajes de error de ICMP reduciendo el MTU del enlace a 1280 bytes, e incluyen una cabecera de fragmentación IPv6 a cada paquete, indicando que este puede ser fragmentado. Esto permite que se realice el proceso de Path MTU Discovery a través del traductor mientras el MTU del enlace sea de 1280 bytes o menor. Cuando el MTU sea menor de 1280 bytes, el nodo enviará paquetes de 1280 bytes que serán fragmentados por enrutadores IPv4 a lo largo del enlace, antes de ser traducidos a IPv4.



✓ Traducción de cabeceras IPv6 a cabeceras IPv4

Los campos de la cabecera IPv4 se traducen como sigue:

Versión	Versión	4
Longitud de la cabecera de internet	Internet Header Length	5 (sin opciones IPv4)
Tipo de Servicio	Type Of Service	Se copia del campo Traffic Class de la cabecera IPv6.
Longitud total	Total Length	Longitud de la carga de datos de la cabecera IPv6, más el tamaño de la cabecera IPv4.
Identificación	Identification	Todos en cero.
Banderas	Flags	La bandera de More Fragments se pone en cero. La bandera de Don't Fragment se pone en uno.
Compensación de fragmento de cabecera	Fragment Offset	Todos en cero
Tiempo de vida	Time to Live	Se copia del campo Hop Limit de la cabecera IPv6. Como el traductor es un enrutador, este debe decrementar el campo Time to Live de la cabecera IPv4 o Hop Limit de la cabecera IPv6 antes de reenviar el paquete. Después de decrementar el valor también debe verificar que este no esté en cero. Si se encuentra en cero, deberá mandar un mensaje de error ("ttl exceded").
Protocolo suma de verificación de cabecera	Protocol Header Checksum	Se copia del campo Next Header de la cabecera IPv6. Se calcula una vez que la cabecera IPv4 ha sido creada.
Dirección de origen	Source Address	Si la dirección origen IPv6 es una dirección de tipo IPv4-traducida entonces se copian los 32 bits más a la derecha. De otra manera, la dirección origen se cambia a 0.0.0.0, evitando que los paquetes sean descartados.
Dirección de destino	Destination Address	Los paquetes IPv6 que son traducidos deben tener una dirección destino IPv6 de tipo IPv4- mapeada. Se copian los 32 bits más a la derecha.

Tabla 25: Traducción de cabeceras IPv6 a cabeceras IPv4

Si el paquete IPv6 contiene una cabecera de fragmentación, entonces los campos se traducen como se indicó anteriormente, con las siguientes excepciones:

Longitud total	Total Length	Longitud de la carga de datos de la cabecera IPv6, menos 8 de la cabecera de fragmentación, más el tamaño de la cabecera IPv4.
Identificación	Identification	Se copian los 16 bits más a la derecha del campo Identification de la cabecera de fragmentación.
Banderas	Flags	



		Se copia la bandera M de la cabecera de fragmentación a la bandera More Fragments. La bandera Don't Fragment se pone en cero, permitiendo que el paquete sea fragmentado por enrutadores IPv4.
Compensación de fragmento de cabecera	Fragment Offset	Se copia del campo Fragment Offset de la cabecera de fragmentación.
Protocolo	Protocol	Se copia del campo Next Header de la cabecera de fragmentación.

Tabla 26: Traducción de cabeceras IPv6 a cabeceras IPv4 (2)

➤ NAT-PT

El método de NAT-PT (Network Address Translator – Protocol Translator) es similar al método de NAT utilizado en IPv4, pero no idéntico. El NAT utilizado en IPv4 consiste en traducir una dirección IPv4 a otra dirección IPv4, mientras que NAT-PT consiste básicamente en la traducción de direcciones IPv4 a direcciones IPv6 y viceversa. NAT-PT utiliza un grupo de direcciones IPv4 (se asume que son únicas globalmente y no privadas) para asignar dinámicamente a los nodos IPv6 cuando estos inicien una sesión para establecer comunicación con algún otro nodo. Todos los paquetes pertenecientes a una misma sesión deberán pasar por el mismo enrutador NAT-PT. Este método utiliza SIIT para la traducción de protocolos, con algunas modificaciones que se explicarán más adelante.

Una parte fundamental de NAT-PT son los ALG's (Application Level Gateways). Estos se utilizan cuando se manejan direcciones IP dentro de la carga de datos del paquete para realizar la traducción, ya que NAT-PT no revisa esa parte del paquete, y por lo tanto no traduce las direcciones en la carga de datos. El ALG más importante es el DNS-ALG, el cual se encarga de mapear las direcciones IPv4 asignadas a un host IPv6.

✓ NAT-PT Tradicional

El NAT-PT tradicional permite a nodos dentro de un dominio IPv6 establecer comunicación con nodos en un dominio IPv4, pero solamente en un sentido, es decir, solamente paquetes al exterior del dominio IPv6. Este se divide en NAT-PT Básico y NAT-PT.

✓ NAT-PT Básico

A continuación se muestra un diagrama para explicar los distintos tipos de NAT-PT:

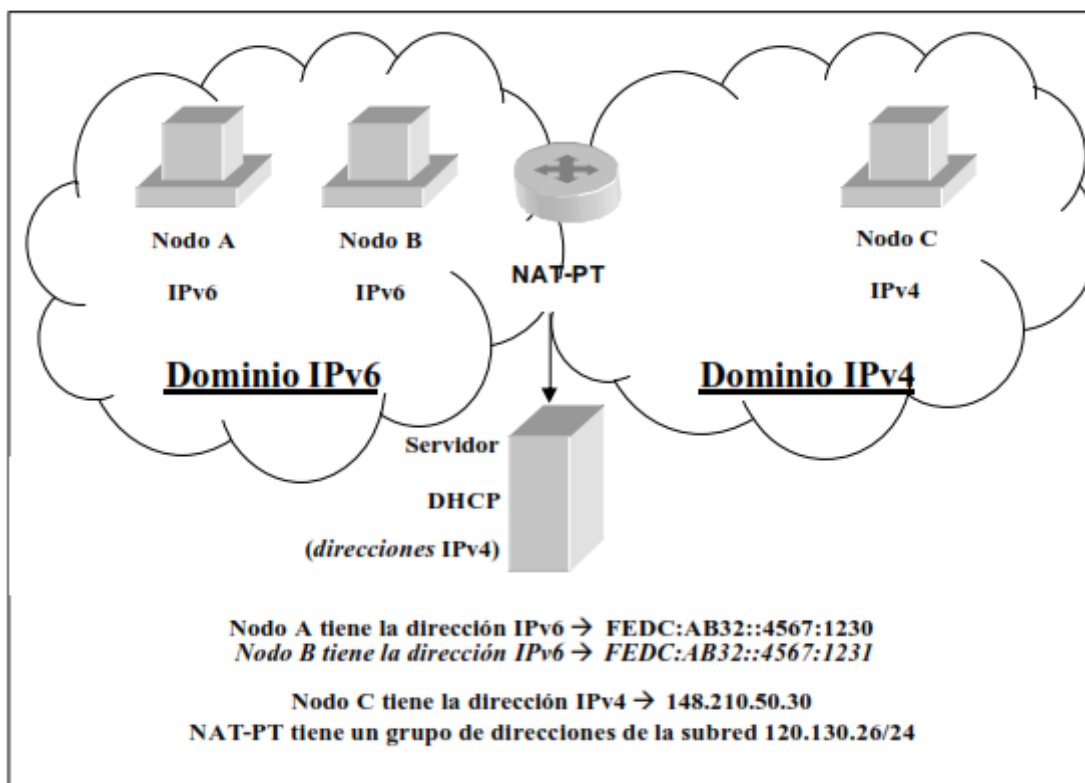


Figura 104: Entorno NAT-PT

Si el grupo de direcciones IPv4 es igual ó mayor que el número de nodos IPv6, entonces se podrá asignar una dirección IPv4 a cada nodo IPv6. Si el grupo de direcciones IPv4 es menor, entonces se tendrán que asignar las direcciones IPv4 dinámicamente.

Supongamos que el nodo A quiere establecer comunicación con el nodo C y crea un paquete con dirección fuente FEDC:AB32::4567:1230 y dirección destino PREFIJO::148.210.50.30. El prefijo PREFIJO::/96 es publicado en el dominio por el NAT-PT, por lo cual todos los paquetes con este prefijo serán direccionados al NAT-PT. Si el paquete no es un inicio de sesión, entonces el NAT-PT deberá conocer el estado de la sesión, así como las direcciones previamente asignadas y el mapeo entre direcciones IPv4 e IPv6. Si el paquete es un inicio de sesión, entonces el NAT-PT le asigna una dirección del grupo de direcciones IPv4 (por ejemplo 120.130.26.1) y traduce el paquete a IPv4. Los parámetros y el mapeo de las direcciones IPv4 e IPv6 se guardan en el NAT-PT el tiempo que dure la sesión. El paquete traducido tiene como dirección fuente 120.130.26.1 y como dirección destino 148.210.50.30. Cualquier paquete de respuesta que sea reconocido como perteneciente a la misma sesión será traducido utilizando la información guardada previamente. La dirección fuente sería PREFIJO::148.210.50.30 y la dirección destino FEDC:AB32::4567:1230.



✓ **NAPT-PT**

NAPT-PT (Network Address Port Translation – Protocol Translation) permite que múltiples nodos IPv6 se comuniquen con nodos IPv4 utilizando una sola dirección IPv4. Esto se logra especificando los puertos TCP/UDP además de la dirección IPv4, lo que permite establecer un gran número de sesiones utilizando una sola dirección IPv4.

Una ventaja de utilizar NAPT-PT por encima de NAT-PT es que soluciona el problema de la escasez de direcciones IPv4. Si se terminaran las direcciones IPv4 utilizadas por NAT-PT, este mecanismo dejaría de funcionar, ya que los nodos IPv6 nuevos no podrían establecer sesiones con redes exteriores.

Supongamos que tenemos NAPT-PT en un **enrutador frontera** (en vez de NAT-PT) y todas las direcciones IPv6 pueden ser mapeadas a una dirección IPv4 única (120.130.26.1). Cuando el nodo A desea establecer una sesión TCP con el nodo C, crea un paquete con dirección fuente FEDC:AB32::4567:1230, puerto TCP fuente 3017, dirección destino PREFIJO::148.210.50.30 y puerto TCP destino 23. Cuando el paquete llega al NAPT-PT, este le asigna una dirección IPv4 respetando el número del puerto de la dirección IPv6, por lo tanto la dirección fuente sería 120.130.26.1, el puerto TCP fuente 1025, la dirección destino 148.210.50.30 y el puerto TCP destino 23.

Cualquier paquete proveniente de la dirección 148.210.50.30 con puerto TCP 23 será reconocido como perteneciente a la misma sesión, y la dirección fuente sería PREFIJO::148.210.50.30, el puerto TCP fuente 23, la dirección destino FEDC:AB32::4567:1230 y el puerto TCP destino 3017.

➤ **NAT-PT Bidireccional**

El NAT-PT Bidireccional, como lo indica su nombre, permite que las sesiones sean iniciadas por hosts dentro del dominio IPv6 o hosts en un dominio IPv4. NAT-PT Bidireccional debe ser utilizado en conjunto con un DNS-ALG para facilitar el mapeo entre direcciones y nombres. Al iniciarse una sesión, la asignación de direcciones IPv4 a paquetes dirigidos al interior o al exterior del dominio IPv6 se manejan de una manera distinta, como se explica a continuación.

✓ **Sesiones al interior del dominio (IPv4 – IPv6)**

Cuando un nodo en el dominio IPv4 envía una requisición de búsqueda de nombre para un nodo dentro del dominio IPv6, esta requisición es enviada al servidor DNS en el dominio IPv6. Como NAT-PT se encuentra en la frontera de ambos dominios, este intercepta la requisición y el DNS-ALG se encarga de traducir la requisición, modificando lo siguiente:

- Cambia el tipo de requisición de A al tipo de requisición AAAA.
- Reemplaza la cadena “IN-ADDR.ARPA” por la cadena “IP6.INT”, así como la dirección IPv4 precedente a la cadena “IN-ADDR.ARPA” con la dirección IPv6 correspondiente (si se ha efectuado previamente un mapeo) en orden inverso.



En caso contrario, cuando se envía una respuesta del servidor DNS del dominio IPv6 al nodo IPv4, NAT-PT intercepta la requisición y el DNS-ALG se encarga de la traducción, modificando lo siguiente:

- Cambia el tipo de respuesta DNS de tipo A al tipo AAAA.
- Reemplaza la dirección IPv6 devuelta por el servidor DNS del dominio IPv6 por la dirección IPv4 asignada previamente por el enrutador NAT-PT. Si la dirección IPv4 no ha sido asignada previamente, se asignará en ese momento.

Supongamos que el nodo C desea establecer una sesión con el nodo A, y solicita una requisición de búsqueda de nombre a su servidor DNS. Este reenvía esta requisición al servidor DNS en el dominio IPv6, pero NAT-PT intercepta esta requisición y solicita al DNS-ALG efectúe la traducción del tipo A al tipo AAAA. Después, es enviada al servidor DNS en el dominio IPv6, quien responde de la siguiente manera:

Nodo A AAAA FEDC:AB32::4567:1230

Cuando se envía esta respuesta al dominio IPv4, NAT-PT intercepta la respuesta y solicita al DNS-ALG efectúe la traducción correspondiente, quedando de la siguiente manera:

NodoA A 120.130.26.1

Esta respuesta de tipo A es enviada al nodo C en el dominio IPv4, y este puede entonces establecer una sesión con el nodo A.

✓ **Sesiones al exterior del dominio (IPv6 – IPv4)**

Los nodos IPv6 pueden obtener las direcciones de nodos IPv4 de un servidor en el dominio IPv4 o del servidor dentro del dominio IPv6. Si el servidor DNS en el dominio IPv6 almacena registros para direcciones de nodos IPv6, así como registros para direcciones de nodos IPv4, entonces las requisiciones hechas por nodos IPv6 para direcciones IPv4 no deberán salir del dominio IPv6 y por lo tanto no serán interceptadas por el NAT-PT. Si el servidor DNS en el dominio IPv6 almacena únicamente registros para direcciones de nodos IPv6, entonces las requisiciones hechas por nodos IPv6 para direcciones IPv4 deberán salir del dominio en busca de un servidor DNS en el dominio IPv4. Esto significa que serán interceptadas por el NAT-PT, y el DNS-ALG se encargará de la traducción.

Supongamos que el nodo A quiere establecer una sesión con el nodo C, por lo que hace una requisición de búsqueda de nombre de tipo AAAA para el nodo C. Como el nodo C puede tener direcciones IPv4 ó IPv6, la requisición se envía sin cambio alguno al servidor DNS en el dominio IPv4, así como una requisición de tipo A.



Si existe un registro de tipo AAAA para el nodo C, este se devuelve al NAT-PT, quien lo envía al nodo A. Si existe un registro de tipo A para el nodo C, este también se devuelve al NAT-PT, entonces el DNS-ALG traduce la respuesta agregando el prefijo correspondiente y lo envía al nodo A, de la siguiente manera:

NodoC A 148.210.50.30

Es traducido a

NodoC AAAA PREF IJO:: 148.210.50.30

El nodo A puede entonces utilizar esta dirección para establecer una sesión con el nodo C.

Traducción de Protocolo

En NAT-PT se utilizan los mismos métodos especificados en SIIT, a excepción de algunas modificaciones, debido a que NAT-PT también traduce direcciones. A continuación se muestran estas modificaciones:

Dirección de origen	Source Address	Los 32 bits más a la derecha son la dirección IPv4. Los 96 bits restantes son el PREFIJO. El prefijo es publicado por el NAT-PT, y todos los paquetes con este prefijo, serán direccionados al NAT-PT.
Dirección de destino	Destination Address	La dirección destino IPv4 es reemplazada por la dirección destino IPv6, basándose en el mapeo previamente establecido.
Dirección de origen	Source Address	La dirección fuente IPv6 es reemplazada por la dirección fuente IPv4, basándose en el mapeo previamente establecido.
Dirección de destino	Destination Address	Los paquetes IPv6 que son traducidos tienen una dirección destino de la forma PREFIJO ::IPv4/96. Los 32 bits más a la derecha de la dirección destino IPv6 se copian la dirección destino IPv4.

Tabla 27: Traducción NAT-PT

➤ BIS

Actualmente existen muy pocas aplicaciones IPv6 en comparación con aplicaciones IPv4. El objetivo en un futuro es que el número de aplicaciones IPv6 sea igual o mayor que el número de aplicaciones IPv4, pero mientras esto ocurre, se necesitan traductores. Un traductor capaz de ejecutar este tipo de traducciones es BIS (Bump-In-the-Stack), que permite a hosts Dual Stack comunicarse con hosts IPv6 utilizando aplicaciones IPv4.

El método de BIS permite a los hosts convertirse en traductores autónomos, sin necesidad de un traductor externo. BIS se encuentra en el área de seguridad IP, y se encarga de verificar los datos que pasan entre TCP/IPv4 y la interfase de red, además de traducirlos a IPv6 y viceversa.

Componentes de BIS

Los hosts Dual Stack necesitan contar con aplicaciones, módulos de TCP/IP y direcciones, tanto para IPv4 como para IPv6. El método de BIS sustituye las aplicaciones IPv6 por tres módulos que le permiten al host



comunicarse con otros hosts utilizando aplicaciones IPv4. Estos tres módulos reciben los nombres de Extension Name Resolver, Address Mapper y Translator.

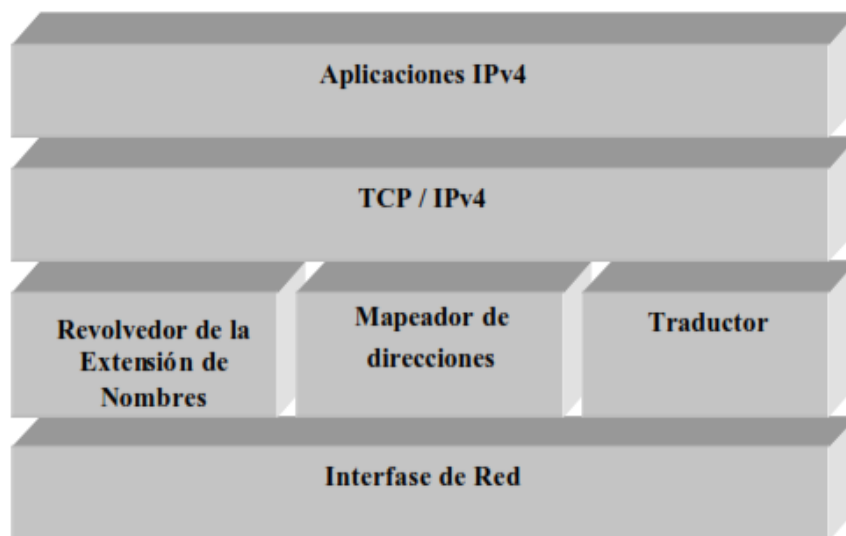


Figura 105: Componentes de BIS.

Resolvidor de la extensión de nombres

Se encarga de responder las requisiciones de nombres de la aplicación IPv4. La aplicación IPv4 hace una requisición de tipo A para obtener la dirección del host con el que quiere establecer comunicación. El módulo de Extension Name Resolver se encarga de crear otra requisición de tipo AAAA y envía ambas requisiciones al servidor DNS. Si este responde a la requisición A, entonces esta es enviada a la aplicación IPv4 y el proceso se lleva a cabo normalmente, es decir, no hay necesidad de una traducción de paquetes. Si el servidor DNS responde a la requisición AAAA, entonces esta es enviada al módulo de Address Mapper para que este le asigne una dirección IPv4 correspondiente y realice el mapeo de las direcciones. Después, el módulo de Extension Name Resolver crea una respuesta a la requisición A con la nueva dirección IPv4 asignada por el Address Mapper y la envía a la aplicación IPv4.

Mapeador de direcciones

Cuenta con un grupo de direcciones IPv4 (pueden ser direcciones privadas). También se encarga de mantener una tabla que consiste en pares de direcciones IPv4 e IPv6. Cuando el módulo de Extension Name Resolver o Translator requieren una dirección IPv4, le notifican a este módulo para que asocie una dirección IPv4 del grupo de direcciones a la dirección IPv6 que se esté utilizando. Esta asociación se debe registrar en la tabla.

Los casos en que se efectúa el registro son los siguientes:

- Cuando el módulo de Extension Name Resolver obtiene respuesta solamente para la requisición de tipo AAAA, y no existe una entrada en la tabla que involucre a la dirección obtenida.



- Cuando el módulo de Translator recibe un paquete IPv6 y no existe una entrada en la tabla que involucre la dirección fuente del paquete.

Existe una excepción en el registro de entradas de la tabla, y esta se presenta al principio de la creación de la tabla, cuando registra un par de direcciones IPv4 e IPv6 propias, de una manera estática.

Traductor

Se encarga de efectuar la traducción entre IPv4 e IPv6, utilizando el método de SIIT. Cuando recibe paquetes IPv4 de aplicaciones IPv4, convierte las cabeceras IPv4 en cabeceras IPv6, y después fragmenta los paquetes IPv6 debido a que las cabeceras IPv6 son dos veces más grandes que las cabeceras IPv4. Cuando recibe paquetes IPv6 de redes IPv6, convierte las cabeceras IPv6 en cabeceras IPv4, pero en este caso no efectúa ninguna fragmentación, ya que no es necesaria.

Comunicación en BIS

En el método de BIS, cuando se desea establecer comunicación entre un host Dual Stack (el cual maneja BIS) y un host IPv6, la comunicación puede ser iniciada por cualquiera de los dos hosts. A continuación se presentan ambos casos.

Comunicación iniciada por host Dual Stack

La aplicación IPv4 hace una requisición de tipo A a su servidor DNS para obtener la dirección del host IPv6. El módulo Extension Name Resolver crea una requisición de tipo AAAA para el host IPv6, y envía ambas requisiciones al servidor DNS. En este caso, como la comunicación está dirigida al host IPv6, solamente se obtendrá respuesta de la requisición tipo AAAA.

El módulo Extension Name Resolver solicita al módulo Address Mapper que le asigne una dirección IPv4 de su grupo de direcciones para asociarla con la dirección IPv6 del host IPv6. El Extension Name Resolver crea una respuesta a la requisición de tipo A, pero con la nueva dirección IPv4 y la envía a la aplicación IPv4. Luego la aplicación envía un paquete IPv4 al host IPv6. El paquete llega al módulo.

Translator, el cual intenta traducir dicho paquete en un paquete IPv6, pero no sabe traducir las direcciones IPv4 fuente y destino. Es entonces cuando el módulo Translator le solicita al módulo Address Mapper que le provea las direcciones IPv6. El Address Mapper revisa su tabla, encuentra las direcciones IPv6 asociadas con las direcciones IPv4 y las envía al módulo Translator. Es entonces cuando este módulo traduce el paquete IPv4 a un paquete IPv6, fragmenta el paquete IPv6 si es necesario y lo envía al host IPv6.



El paquete IPv6 llega al host IPv6 el cual responde enviando paquetes IPv6 al host Dual Stack, que son interceptados por el módulo Traductor de este mismo. El Traductor obtiene las direcciones IPv4 asociadas con las direcciones IPv6, traduce el paquete IPv6 a un paquete IPv4 y lo envía a la aplicación IPv4.

Comunicación iniciada por host IPv6

El host IPv6 hace una requisición de tipo AAAA a su servidor DNS para el host Dual Stack, y envía un paquete IPv6 a la dirección IPv6 obtenida. El paquete IPv6 es interceptado por el módulo Traductor del host Dual Stack.

El módulo Traductor intenta traducir el paquete IPv6 en un paquete IPv4, pero no sabe traducir las direcciones IPv6 fuente y destino. Es entonces cuando el módulo Traductor le solicita al módulo Address Mapper que le provea las direcciones IPv4.

El módulo Address Mapper revisa su tabla y encuentra únicamente la dirección IPv6 destino, la cual fue registrada en la creación de la tabla, pero no encuentra la dirección IPv6 fuente. Después, el módulo Address Mapper asocia una dirección IPv4 de su grupo a la dirección IPv6 fuente, y envía las direcciones IPv4 fuente y destino al módulo Traductor. Es entonces cuando este módulo traduce el paquete IPv6 a un paquete IPv4 y lo envía a la aplicación IPv4 del host Dual Stack. La aplicación le envía un paquete IPv4 al host IPv6 como respuesta y sigue el proceso descrito en la sección anterior.

➤ NAT64 (Network Address Translation IPv6 a IPv4)

Es un mecanismo que permite a un hosts que solamente tienen conectividad IPv6 comunicarse con un hosts que tiene solamente conectividad con IPv4. Cuenta principalmente con dos mecanismos, el de traducción de direcciones y el de traducción de protocolos, lo que permite a una red IPv6 acceder a servicio de red IPv4, sin en el requerimiento de una red doble pila.

En una red solo IPv6, las direcciones IPv4 se mapean dentro de un prefijo IPv6 el cual debe tener suficientes bits de host para mapear todo el espacio IPv4. Normalmente se utiliza el prefijo 64::ff9b/96 (llamado Prefijo Well-Known).

En caso, que los hosts solo IPv6 puedan comunicarse con los hosts solo IPv4 hace falta una traducción a nivel de DNS. Este es el rol del DNS64, el cual se encarga de recibir las consultas DNS de los hosts solo IPv6 y modificar las respuestas de tal manera de incluir registros AAAA que mapean las direcciones IPv4 dentro del prefijo NAT64.

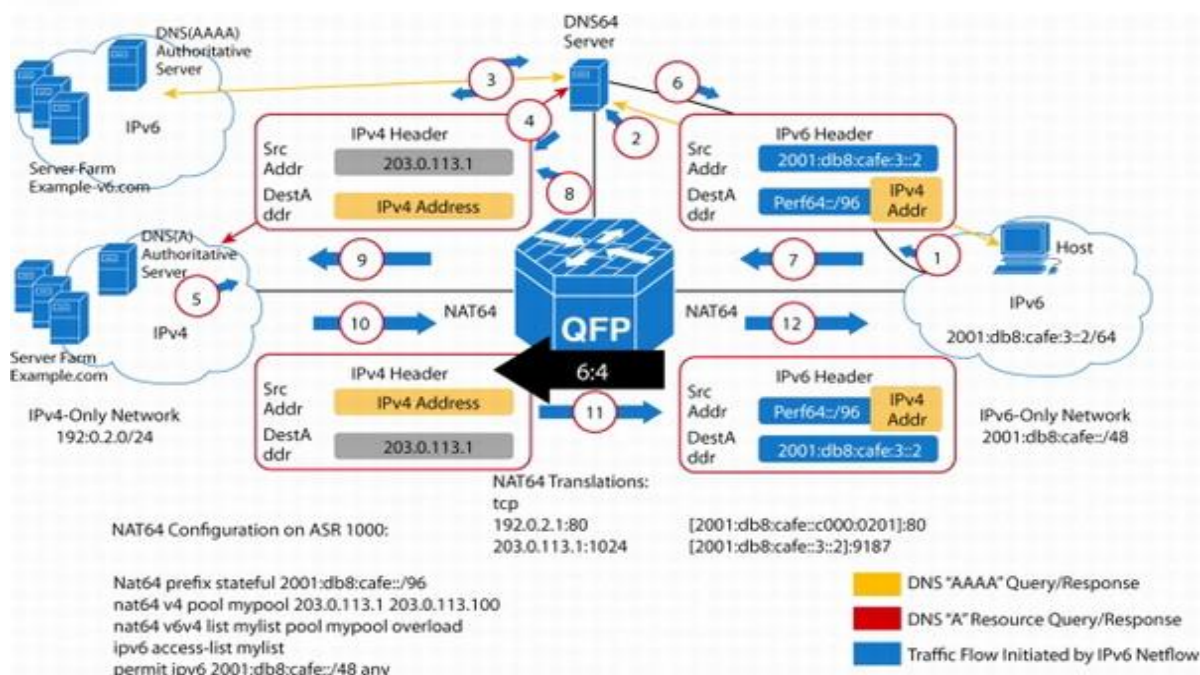


Figura 106: Representación de NAT-64

Un entorno de NAT64 puede verse como un dispositivo de red (un Router) con al menos dos interfaces. Uno de los interfaces está conectado a la red IPv4, y el otro a la red IPv6. La red estará configurada de modo que los paquetes de la red IPv6 a la red IPv4 son encaminados a través de este router. El router realizará todas las traducciones necesarias para transferir paquetes de la red IPv6 a la red IPv4, y viceversa.

La traducción no es simétrica, dado que el espacio de direcciones IPv6 es mucho mayor que el de direcciones IPv4 (compara: 2128 en IPv6 y 232 en IPv4), por lo que no es posible una traducción una-una.

Para poder llevar a cabo la traducción, el equipo NAT64 debe mantener un mapeo de direcciones IPv6 a IPv4 (es decir, mantiene estado). Este tipo de mapeo de direcciones se configura estáticamente por los administradores del sistema o, habitualmente, se crea automáticamente cuando llega el primer paquete IPv6 al servidor NAT64. Después de que se haya creado este flujo, los paquetes pueden pasar en ambas direcciones. En general, NAT64 está diseñado para usarse cuando las comunicaciones son iniciadas por los hosts IPv6.

Traducción sin estado NAT64-sin estado entre IPv4 e IPv6

RFC6145 (/ ICMP Traducción Algoritmo IP) sustituye RFC2765 (sin estado IP / ICMP Traducción Algoritmo (SIIT)) y proporciona un mecanismo sin estado para traducir un encabezado IPv4 en una cabecera IPv6 y viceversa.



La clave para la traducción sin estado está en el hecho de que la dirección IPv4 está incrustado directamente en la dirección IPv6. Una limitación de la traducción sin estado NAT64 es que se traduce directamente sólo a las opciones de IPv4 que tienen contrapartes directas a IPv6, y que no se traduce cualquier cabecera de extensión IPv6 más allá de la cabecera de extensión de fragmentación.

Con un NAT64, un rango específico de direcciones IPv6 representará sistemas IPv4 dentro del mundo IPv6. Esta gama se debe configurar manualmente en el dispositivo de traducción. Dentro del mundo IPv4 todos los sistemas IPv6 han correlacionado directamente las direcciones IPv4 que se pueden algorítmicamente asignar a un subconjunto de las direcciones IPv4 del proveedor de servicios. Por medio de este algoritmo de asignación directa no hay necesidad de mantener el estado de cualquier ranura de traducción entre IPv4 e IPv6. Este algoritmo de asignación requiere que los hosts de IPv6 pueden asignar direcciones IPv6 específicas, utilizando la configuración manual o DHCPv6.

Dirección de Stateful NAT64 de la red y el Protocolo traducción del Clientes IPv6 a servidores IPv4.

Múltiples dispositivos IPv6 en una sola dirección IPv4. Se puede suponer que esta tecnología se utiliza principalmente donde sólo IPv6 redes y clientes (es decir. Los teléfonos móviles, IPv6 única inalámbricos, etc) necesitan tener acceso al Internet IPv4 y sus servicios.

La gran diferencia con NAT64 con estado es la eliminación de la unión entre la dirección IPv6 y la dirección IPv4 algorítmica. A cambio, se crea el estado en el dispositivo NAT64 para cada flujo. Además, NAT64 sólo admite los flujos de IPv6-iniciado. A diferencia de apátridas NAT64, NAT64 stateful no 'no 'consumen una única dirección IPv4 para cada dispositivo IPv6 que quiere comunicar a la Internet IPv4. Más prácticamente, esto significa que muchos sólo IPv6 usuarios consumen solamente única dirección IPv4 en forma similar a la de direcciones de red IPv4 a IPv4 y traducción puerto de obras. Esto funciona muy bien si se inicia la petición de conectividad de las IPv6 hacia la Internet IPv4. Si un sólo IPv4 dispositivo quiere hablar con un sólo IPv6 del servidor, por ejemplo, se requiere la configuración manual de la ranura de traducción, por lo que este mecanismo sea menos atractivo para proporcionar servicios IPv6 hacia Internet IPv4.

Diferencias ente Stateles NAT64 y con Statefull Nat64

NAT64 sin estado	NAT64 con estado
1: 1 traducción	1: traducción N
No conservación de direcciones IPv4	Conserva dirección IPv4



Asegura la transparencia y la escalabilidad de extremo a extremo de dirección	Utiliza dirección de sobrecarga, por lo tanto, carece de transparencia en la dirección de extremo a extremo
Ningún Estado o consolidaciones creado en la traducción	Estado o fijaciones se crean en cada traducción única
Requiere IPv4-traducible direcciones IPv6 asignación (requisito obligatorio)	No hay requisito de la naturaleza de la asignación de direcciones IPv6
Requiere ya sea manual o asignación de direcciones basado DHCPv6 para hosts de IPv6	Libre de elegir cualquier modo de asignación de direcciones IPv6 a saber. Manual, DHCPv6, SLAAC

Tabla 28: Diferencia de NAT64 con estado y sin estado.

DNS64

DNS64 es una pasarela de nivel de aplicación para el protocolo DNS que genera registros de tipo AAAA (AAAA RRs) a partir de registros A (A RRs). DNS64 permite a equipos que utilizan únicamente IPv6 utilizar nombres de dominio de equipos IPv4 para iniciar la comunicación.

cuando un servidor DNS64 recibe una petición por un registro AAAA generada por un iniciador IPv6 ,busca un registro tipo AAAA. Si el registro AAAA no existe para el nodo contactado (que es el caso habitual con nodos que utilicen sólo IPv4), el DNS64 realiza una búsqueda del registro de recurso tipo A. Si un registro “A” es encontrado, DNS64 crea un registro AAAA sintético añadiendo el prefijo “Pref64::/n” del NAT64 a la dirección IPv4 del nodo a contactar (y en caso de que “n” es menor que 96, adicionalmente un sufijo, al igual que para NAT64). El registro AAAA sintético es devuelto al nodo IPv6 origen, el cual inicia una comunicación IPv6 con la dirección IPv6 asociada a la dirección IPv4 destino.

El paquete es enrutado hacia el NAT64 local, el cual realiza el mapeo de direcciones, tanto origen como destino, IPv6 a IPv4 descrito anteriormente. Es importante observar que el DNS64 y el NAT64, no comparten ninguna información de estado. En particular, cuando el DNS64 genera respuesta sintética, no se guarda ningún tipo de estado en el NAT64.

La única información compartida en ambos es el prefijo Pref64::/n, que es definido por dominio.

Por defecto, ambos NAT64 y DNS64 utiliza el prefijo **“Well-Known”** mencionado con anterioridad, por que no es necesaria una configuración manual en ninguno de los dos.

Recorrido por NAT64/DNS64

Para este recorrido consideramos la topología que se muestra a continuación. El NAT64 utiliza el prefijo Well-Known 64:ff9b::/96 para mapear direcciones IPv4 a IPv6 y tiene asignada la dirección **T** a su interfaz IPv4 hacia la red exterior. El servidor DNS local implementa la funcionalidad DNS64 y utiliza el prefijo Well-Known para la síntesis de los registros AAAA. Los equipos IPv6 realizan búsquedas recursivas contra el servidor local de DNS.

A continuación describimos como H1 inicia la comunicación con H2:

- 1) H1 realiza una búsqueda DNS para la dirección IPv6 de H2 enviando una petición DNS de un registro AAAA al servidor DNS/DNS64 local.
- 2) El servidor local de DNS/DNS64 resuelve la petición, y descubre que no existe registro tipo AAAA para H2.

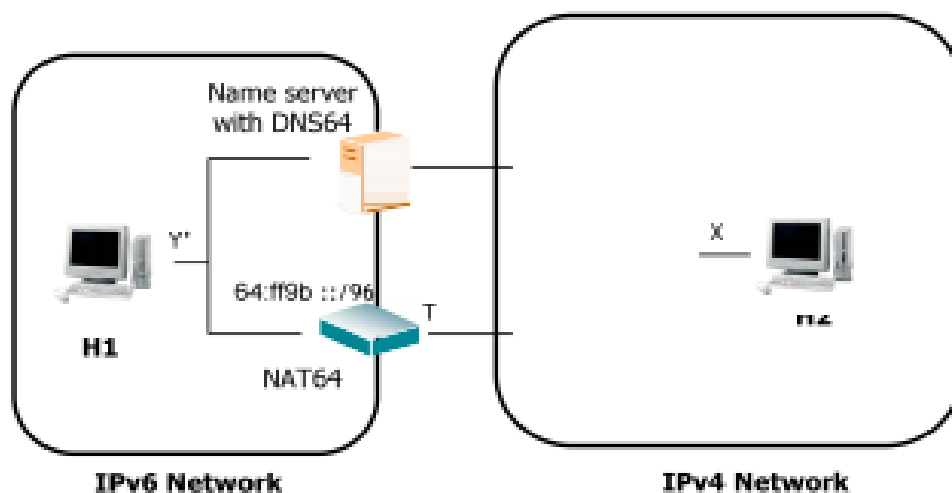


Figura 107: Función de DNS64

- 3) El servidor DNS/DNS64 busca un registro A para H2, obteniendo la dirección IPv4 X.
- 4) El servidor DNS/DNS64 sintetiza un registro AAAA añadiendo el prefijo 64:ff9b::/96 a la dirección IPv4 X, e incluye la dirección en la respuesta a H1.
- 5) Después de recibir el registro sintético AAAA, H1 manda un paquete hacia H2 desde la dirección de transporte origen (Y', y) a la dirección de transporte destino (64:ff9b:X, x), donde y y x son puertos elegidos por H2.



- 6) El paquete se enruta hacia el interfaz IPv6 del NAT64 (dado que 64:ff9b::/96 ha sido asociado a esta interfaz), y el NAT64 realiza las siguientes acciones:
 - Selecciona un puerto libre t y crea una entrada $(Y', y) \leftrightarrow (T, t)$
 - Traduce la cabecera IPv6 en la cabecera IPv4 utilizando traducción sin estado.
 - Incluye en el paquete (T, t) como dirección de transporte origen y (X, x) como dirección de transporte destino.
 - El NAT64 envía el paquete a la red IPv4.
- 7) El nodo H2 recibe el paquete y contesta enviando un paquete con destino la dirección de transporte (T, t) y la dirección de transporte (X, x) como origen.
- 8) El paquete se enruta hacia el NAT64 a través de la red IPv4. El NAT64 busca la entrada conteniendo (T, t) . Cuando la entrada se encuentra:
 - El NAT64 traduce el paquete IPv4 a un paquete IPv6 usando la traducción sin estado.
 - El NAT64 incluye en el paquete la dirección de transporte (Y', y) como dirección de origen y $(\text{Pref64}:X, x)$ como dirección transporte destino.



16. Movilidad en IPv6

Se entiende por movilidad a la capacidad que tiene un nodo de una red, para mantener la misma dirección IP, a pesar que se desplace físicamente a otra red. Es decir que sin importar su ubicación este puede seguir siendo accesible a través de su misma dirección IP.

Sin esta capacidad, los paquetes destinados a un nodo móvil, no podrán llegar a destino mientras dicho nodo, se encuentre alejado de su enlace principal.

Para que un nodo tenga la capacidad de movilidad, la misma debe ser habilitada en el mismo. Mientras este nodo se encuentra en su red, Home Network, la dirección IP que tiene asignada se conoce como Home Address. Siempre que su ubicación sea en su red origen, los paquetes enviados a esa dirección serán ruteados utilizando los mecanismos tradicionales de Internet.

Cuando el nodo se desplaza hacia otra red, adquiere una nueva dirección, conocida como Careof Address, con igual prefijo de la red visitada.

Una vez que configuró su nueva dirección, debe informársela a un nodo ubicado en su Home Network, que se conoce como Home Agent. Este proceso de asociar la home address con la nueva care-of address se conoce como Binding. El momento en que el nodo móvil se mueve a otra red es el punto crítico del proceso. Esto se conoce como handover, y es el momento en el que el nodo móvil pierde conectividad con el otro extremo hasta que termine todo el procedimiento de obtener la nueva dirección y registrarse con el home agent. Este lapso debe ocupar el menor tiempo posible para evitar que se pierdan muchos paquetes, que luego tendrán que ser retransmitidos, porque mientras se encuentra en este estado el nodo no es capaz de recibir paquetes enviados a su home address.

El handover también se produce cuando un nodo se mueve entre diferentes Access Points (AP) pertenecientes a una misma red wireless, es decir, al desplazarse el nodo va cambiando su asociación entre los diferentes AP pero mantiene su dirección IP. Este proceso se realiza a nivel de enlace, y es transparente a las capas superiores. Estas no se enteran de este desplazamiento, con lo cual la dirección IP no se modifica.

El diagrama ilustra el problema de movilidad en un entorno de Internet. Se muestra un **Dominio local** con un **Router (dominio local)** y un **RAL de origen**. Un **MN (Mobile Node)** se mueve desde el dominio local hacia un **Dominio remoto** que contiene un **Router (dominio remoto)** y un **RAL remota**. El **MN** ahora está etiquetado como **NODO MÓVIL**. El **Internet** conecta los dos dominios. El **CN (Correspondent Node)** en el dominio remoto recibe tráfico desde el **RAL de origen** a través del Internet, pero no sabe que el **MN** se ha movido, lo que causa problemas de entrega.

Objetivo: MN pueda conservar permanentemente su dirección IP de origen o nativa, independientemente de su ubicación física, y seguir manteniendo sus comunicaciones y recursos como si estuviera en su **RAL** de origen.

- **Nodo móvil (MN)**
Terminal con funcionalidad IP móvil que se conecta a otra red IP manteniendo su dirección IP
- **Agente de casa (HA)**
Router con funcionalidad IP móvil en la red original de MN.
Representante de MN en su ausencia y una vez MN haya registrado, previamente, en HA su CoA.
- **Nodo corresponsal (CN)**
Terminal con funcionalidad IP móvil que está manteniendo una comunicación con MN.
- **Dirección Care-of-Address (CoA)**
Dirección IP temporal de MN en la RAL destino.
La dirección CoA se utiliza exclusivamente para definir el extremo del túnel IP móvil en MIPv6.

16.2 Túnel IP Móvil en la movilidad IPv6

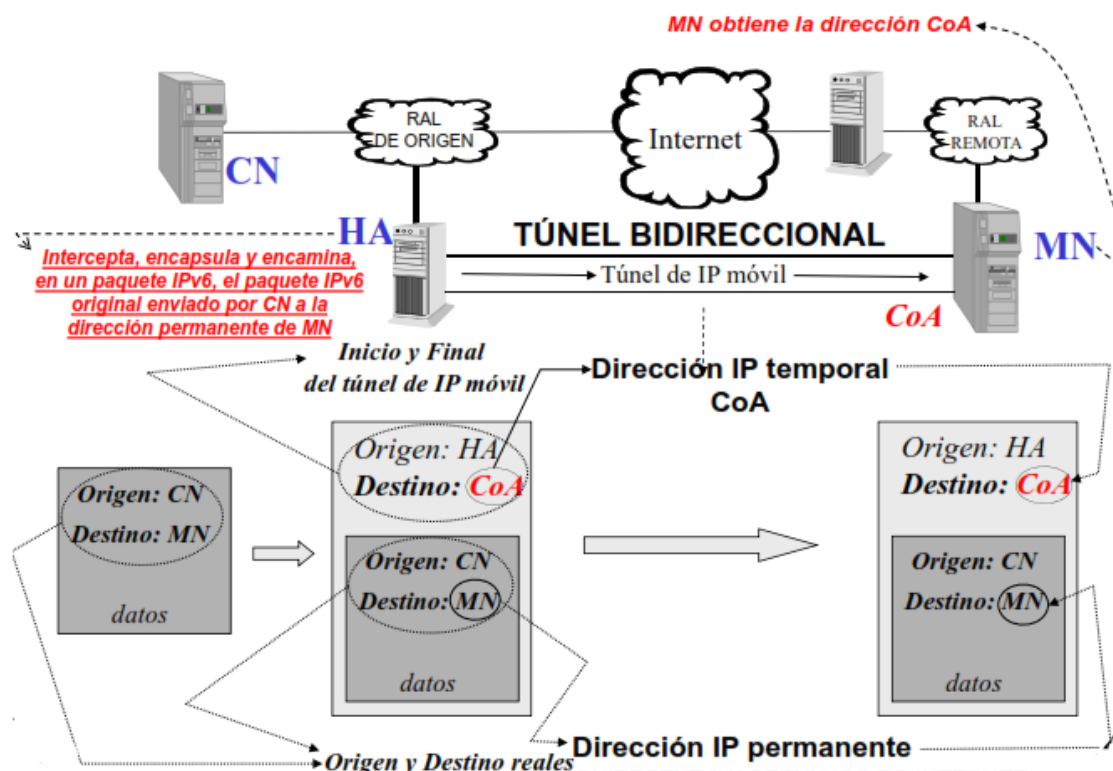


Figura 109: Túnel IP móvil en MIPv6.

Cuatro Etapas en la Movilidad IPv6

1. **Descubrimiento de Agentes (Agent Discovery):** Proceso en MN de detección del Router local o RAL local o Router remoto o RAL remota.
2. **Adquisición de una CoA (New-CoA adquisición):** Proceso de obtención de una nueva CoA por MN.
El nodo móvil (MN) cuando cambia de red tiene dos direcciones:
Dirección permanente de la RAL origen o “de casa” (Home Address).
Dirección temporal de la RAL destino (Care-of-Address).
3. **Registro (Registration) con HA:** Actualización de la tabla de asociación (binding) del HA con un paquete IPv6, procedente de MN, con la cabecera de extensión de movilidad y la opción de actualización de la Asociación (Binding Update):
Dirección permanente-CoA.
4. **Triángulo de encaminamiento y encapsulación (routing and tunneling):** Túnel IP móvil a través de HA permitiendo comunicaciones entre MN y uno o varios CN.

El nodo correspondiente (CN) no se da cuenta de que el terminal móvil está en otra red y envía paquetes a la RAL origen de MN.

Los paquetes con destino al MN son interceptados por HA (Home Agent) y enviados por un túnel de IP móvil a la red remota donde está MN.

Las respuestas de MN a CN y, posteriormente, de CN a MN por él o varios CN triángulo con HA. MN envía los paquetes con su CoA (dirección origen).

Si pusiera su dirección permanente, el router podría (por seguridad) hacer un filtrado y descartar el paquete, porque la dirección de origen no coincide con el prefijo de red. (Evitando un posible ataque a una máquina por Internet desde una máquina con una dirección “extraña” en dicha red destino).

➤ Optimización del triángulo o respuestas directas:

Las respuestas de MN a CN y, posteriormente, de CN a MN, directamente entre ellos, es decir, rompiendo el triángulo y, por tanto, sin pasar por HA.

Tres Actores, 4 etapas y 2 Tipos de Comunicaciones

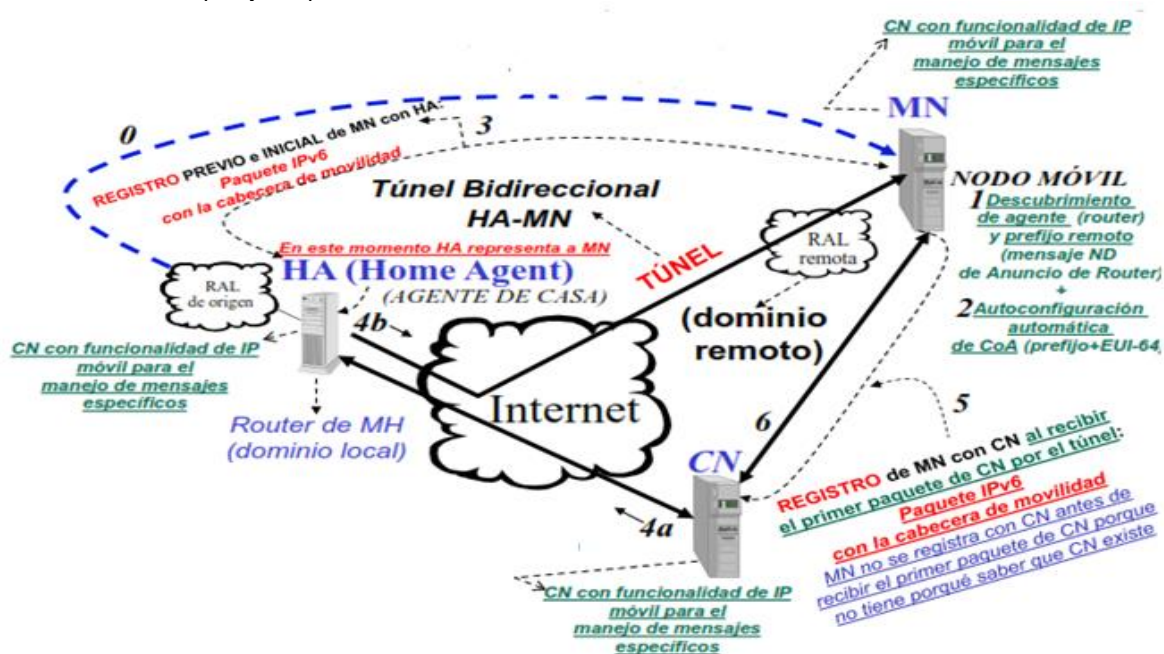


Figura 110: Optimización del triángulo. O repuestas directas.



➤ **Dos comunicaciones posibles:**

a) Túnel Bidireccional (triángulo) de CN a MN y de MN a CN siempre vía HA.

Cuando un CN quiere contactar con el MN, lo que hace es intentar contactar con el MN a través de su HoA ya que es la dirección fija conocida por el CN. Los paquetes enviados a la red del operador y dirigidos a la HoA del MN, son interceptado por el HA, encapsulados en un paquete MIPv6 y redirigidos hacia la nueva dirección CoA que el nodo móvil tiene en la red visitada.

El MN contesta al CN encapsulando los paquetes de datos en un paquete MIPv6 y se lo envía al HA, que extrae el paquete original y se lo envía al CN.

b) Respuestas directas (sin HA) de MN a CN y, luego de CN a MN.

Si el CN no tiene soporte MIPv6, es posible que el MN contacte con el CN para informarle de que un IPv6 cuando está en la red visitada es la CoA y no la HoA, de forma que el CN envía los paquetes de datos directamente a la CoA del nodo móvil. Este procedimiento se denomina **Route Optimization** y es una mejora en el camino seguido por los paquetes, ya que no tiene que pasar por el HA evitando retrasos innecesarios. Si el CN no posee soporte MIPv6 no es posible que el CN y el MN se comuniquen usando Route Optimization.

Proceso de Señalización o Envío de Mensajes entre las 3 entidades funcionales en las 2 primeras etapas de movilidad.

1. Descubrimiento de Agentes o detección del router local en la RAL local o router remoto en la RAL remota.

Mensaje ND de Anuncio de Router: Transmitido regularmente por un router local a MN para indicar su existencia y características mediante un mensaje ICMPv6 de anuncio de Router (Tipo = 134) con una o más opciones de información del mensaje ND (dirección MAC del router, prefijo de red, MTU, información del HA, etc.) con información específica para los MN.

Por el prefijo de red, MN sabe si está en “casa” o “fuera”.

2. Adquisición de una CoA: Sólo cuando MN detecta que está en una RAL remota obtiene CoA por autoconfiguración automática (prefijo de red + EUI-64) o por DHCPv6.

Formato del Mensaje ND de Anuncio de Router

Etapas de Descubrimiento de Agentes

Formado por un Mensaje ICMPv6 de Anuncio de Router (134)

Los mensajes del protocolo ND se construyen con mensajes ICMPv6

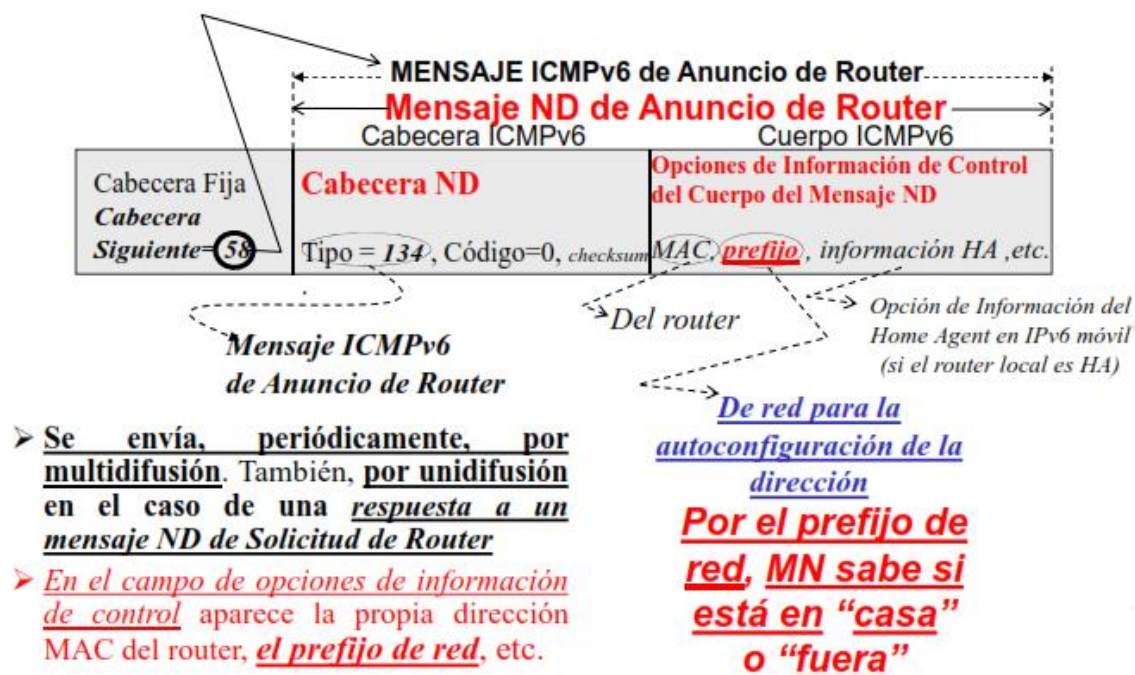


Figura 111: Formato de mensaje .ND de anuncio de router.

Proceso de Señalización o Envío de Mensajes entre las 3 Entidades Funcionales en la 3ª Etapa de Movilidad de IPv6.

1. **Registro con HA:** MN envía a su HA un paquete IPv6 con la cabecera de extensión de movilidad que contiene la opción de Actualización de la Asociación (Binding Update):

Dirección permanente-CoA para la actualización de la tabla de asociación de HA

HA comprueba si MN está en la RAL de origen mediante un mensaje ND de solicitud de vecino creado a través de un mensaje ICMPv6 de solicitud de vecino (tipo 135) con la opción dirección IPv6 permanente de MN: La respuesta es un mensaje ICMPv6 de anuncio de vecino (tipo 136) que incluye la dirección MAC del vecino MN.

No hay respuesta de MN si MN no está en su RAL de origen.

Cabecera de Extensión de Movilidad.

Opción de actualización de la asociación (binding update): Dirección permanente de MN y CoA registro inicial de MN con HA cuando MN obtiene su CoA en la RAL remota y, posterior, registro de MN con CN.



Una vez realizado el registro de MN en el router HA, MN realiza también el registro en CN y, a continuación, se lleva a cabo con éxito la prueba de encaminamiento.

– Por tanto, es posible enviar directamente paquetes de MN a CN y de CN a MN sin pasar por el router HA.

Tras recibir un primer paquete de CN por el túnel vía HA, MN transmite directamente hacia CN.

Cuando MN regresa a casa (RAL de origen) también envía al HA una Cabecera de Extensión de movilidad con la opción correspondiente para indicarle de su regreso y de la no necesidad de una CoA.

Cualquier mensaje que incluya una actualización de la asociación (binding update) debe incluir una cabecera AH y otra ESP (o sólo ESP con funcionalidad AH).

La movilidad IPv6 puede hacer uso de IPSec para todos los requerimientos de seguridad, como la autenticación, integridad y confidencialidad.

16.3 Nueva Cabecera de Extensión IPv6

Cabecera de extensión de movilidad

Cabecera siguiente en la anterior cabecera = 135.

Para transportar mensajes de movilidad en IPv6 (Mobile IPv6).

- Permite a MN registrar y asociar su CoA con HA para túneles bidireccionales.
- Permite a MN registrar y asociar su CoA, directamente, con CN para comunicaciones directas sin pasar por HA.

Tiene que ser la última cabecera en el paquete IPv6.

- Cabecera siguiente = 59, (significa que no hay más cabeceras) y sin PDU del nivel superior.

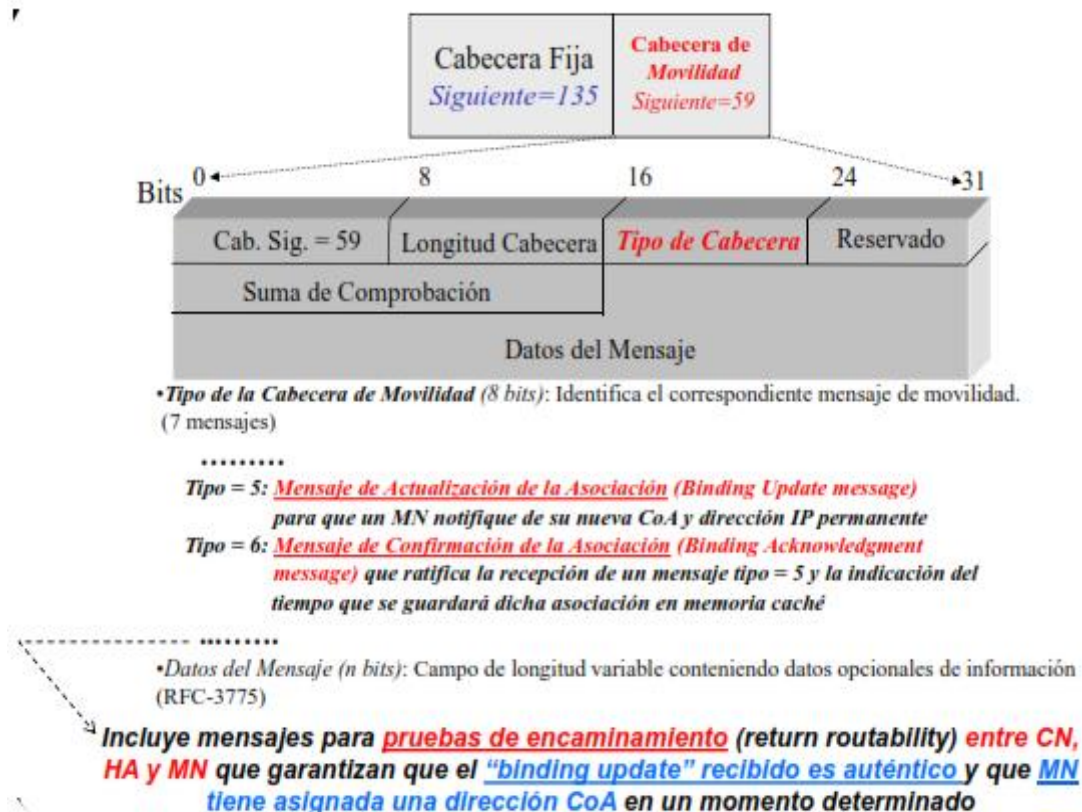


Figura 112: Cabecera de extensión de IPv6.

Proceso de Señalización o Envío de Mensajes entre las 3 Entidades Funcionales en la 4ª Etapa de la Movilidad IPv6.

3. Encaminamiento y encapsulación

HA se ocupa de interceptar cualquier paquete destinado a la dirección permanente de MN.

Encaminamiento en triángulo: Si el CN se comunica con un MN, los paquetes se encaminan desde el CN hacia el HA que los encapsula, y encamina al extremo del túnel con MN y cuya terminación es CoA.

➤ OPTIMIZACIÓN (evitar el triángulo pasando por HA) o envío directo de respuestas entre MN y CN:

REGISTRO: Una vez MN recibe por el túnel paquetes de uno o más CN, envía hacia dichos CN que se están comunicando con él, paquetes con la cabecera de extensión de movilidad conteniendo la opción de actualización de la asociación (Binding Update).



- Datos: Como MN puede responder directamente a CN sin pasar por el túnel con HA, pone su dirección temporal (CoA) como dirección origen del paquete y su dirección permanente en la cabecera de extensión de opciones para el destino.
- De esta forma, la dirección IP CoA es transparente al nivel de red, transporte y aplicación de CN.
- Posteriormente, CN envía paquetes IPv6 con una cabecera de extensión de encaminamiento que contiene la dirección IP permanente de MN.

Cabecera de Extensión de Opciones para el Destino

Por ser una respuesta para el destino de la solicitud previa.

Tráfico de MN a CN (directamente sin pasar por HA).

Los paquetes emitidos por MN llevan la dirección origen CoA en la cabecera fija y la dirección permanente de MN en la opción para que la dirección CoA sea transparente para el nivel de red, transporte y aplicación de CN.

- La acción que debe realizar CN, al recibir el paquete, es cambiar la dirección de origen CoA por la dirección permanente de CN que está en la cabecera de extensión de opciones para el destino.

Cabecera de Extensión de Encaminamiento (tipo 2)

Por ser una solicitud de servicio que se ha de encaminar hacia un destino temporal (CoA).

Tráfico de CN a MN (directamente sin pasar por HA).

Los paquetes emitidos por CN llevan la dirección destino CoA en la cabecera fija y la dirección permanente del destino (MN) en la opción para que la dirección CoA sea transparente para el nivel de red, transporte y aplicación de MN.

- La acción que debe realizar MN, al recibir el paquete, es cambiar la dirección destino CoA por la dirección permanente de CN que está en la cabecera de extensión de encaminamiento.



Cabeceras de Extensión en Movilidad IPv6

Valor decimal de cabecera siguiente	Cabecera de extensión	
0	Cabecera de extensión de opciones salto a salto	
43	Cabecera de extensión de encaminamiento	RFC-2460
44	Cabecera de extensión de fragmentación	
51	Cabecera de extensión AH	RFC-4302
50	Cabecera de extensión ESP	RFC-4303
60	Cabecera de extensión de opciones para el destino	RFC-2460
135	Cabecera de movilidad	RFC-3775

Tabla 29: Cabeceras de extensión en MIPv6

Cabecera de Extensión de Movilidad IPv6 de Opciones para el Destino *Tráfico de MN a CN (directamente sin pasar por HA)*

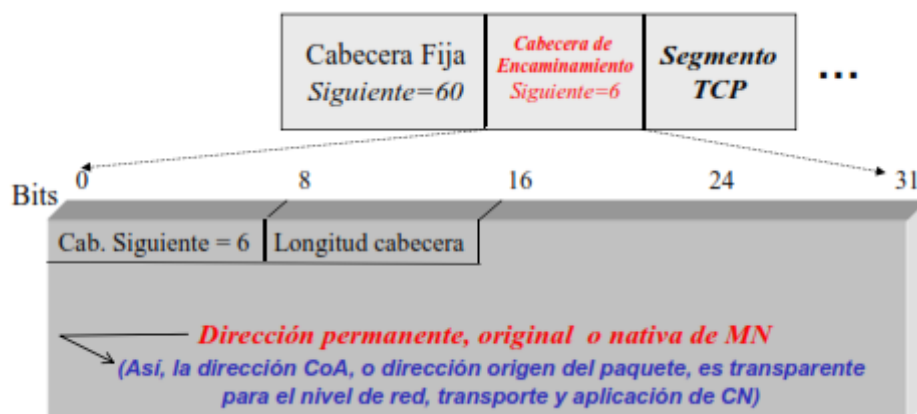


Figura 113: Tráfico de MN a CN.

Cabecera siguiente (8 bits).

Longitud cabecera (8 bits): Longitud de la cabecera en bloques de 8 octetos sin incluir los primeros 8 octetos.

Dirección permanente, original o nativa del emisor.

Cabecera de Extensión de Movilidad IPv6 de Encaminamiento Tipo 2 Tráfico de CN a MN (directamente sin pasar por HA)

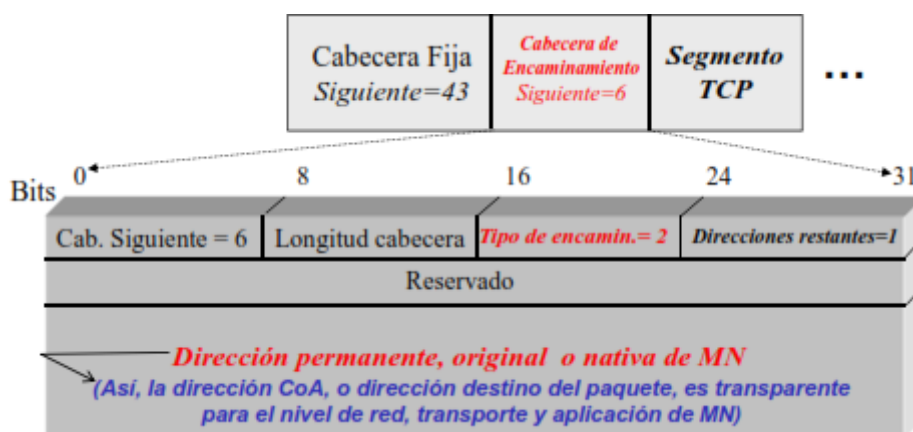


Figura 114: Cabecera de extensión MIPv6 tráfico de CN a MN.

- Cabecera siguiente (8 bits).
- Longitud cabecera (8 bits): Longitud de la cabecera en bloques de 8 octetos sin incluir los primeros 8 octetos.
- Tipo de encaminamiento (8 bits): Actualmente se ha definido el tipo cero (encaminamiento estricto y no estricto de IPv4) y tipo 2 para IP móvil.
- Direcciones restantes (8 bits): Número de destinos intermedios (encaminamiento tipo cero) o 1 (encaminamiento tipo 2).
- Reservado (8 bits): A ceros.
- Dirección permanente, original o nativa del destinatario.

16.4 Ejercicio de análisis de movilidad

Caso de estudio

Una organización tiene oficinas en Madrid, León, Valencia y Barcelona. En cada oficina los equipos informáticos están conectados mediante tecnología Ethernet y utilizan un único router para su conexión a Internet.

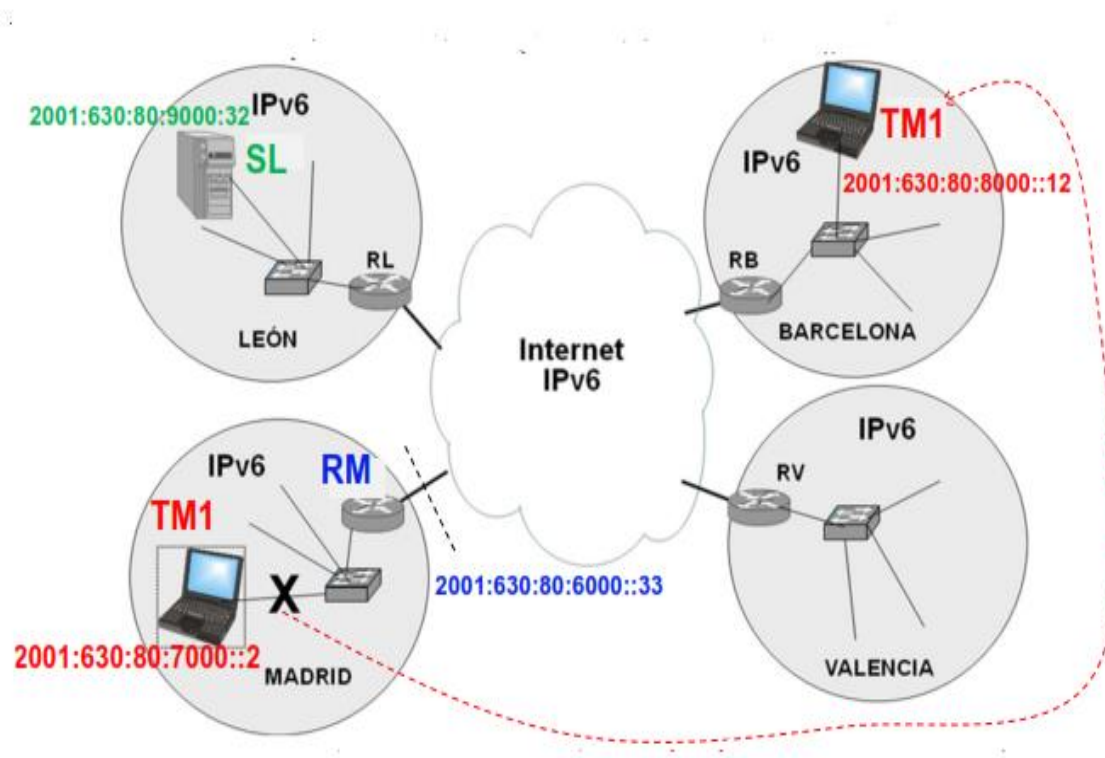


Figura 115: Caso de estudio de MIPv6.

Un empleado de la oficina de Madrid que utiliza un terminal portátil TM1 se traslada a trabajar durante un mes a la oficina de Barcelona; por lo que ahora conecta su terminal TM1 a la red de dicha oficina de Barcelona. Se desea que siga disponiendo de la misma conectividad IPv6 que tenía en la oficina de Madrid, es decir, que no se modifiquen, en los correspondientes servidores DNS, los registros de la dirección IP del terminal TM1.



Incisos a resolver:

Inciso 1.

Explique la funcionalidad que deben tener los routers RM, RB y el terminal TM1; así como las acciones concretas que realizan cada uno de ellos.

Inciso 2.

Una vez realizado el registro de TM1 en el router RM y posteriormente, en el servidor SL, se envían directamente paquetes (sin pasar por el router RM) de SL a TM1 a SL.

- Indicar la estructura de la cabecera IP de un paquete que envía el servidor SL al terminal TM1 (localizado en la oficina de Barcelona), detallando el contenido de los campos que se conozcan.

Inciso 3.

Indicar la estructura de la cabecera IP de un paquete que envía el terminal TM1 (localizado en la oficina de Barcelona) al servidor SL, detallando el contenido de los campos que se conozcan.

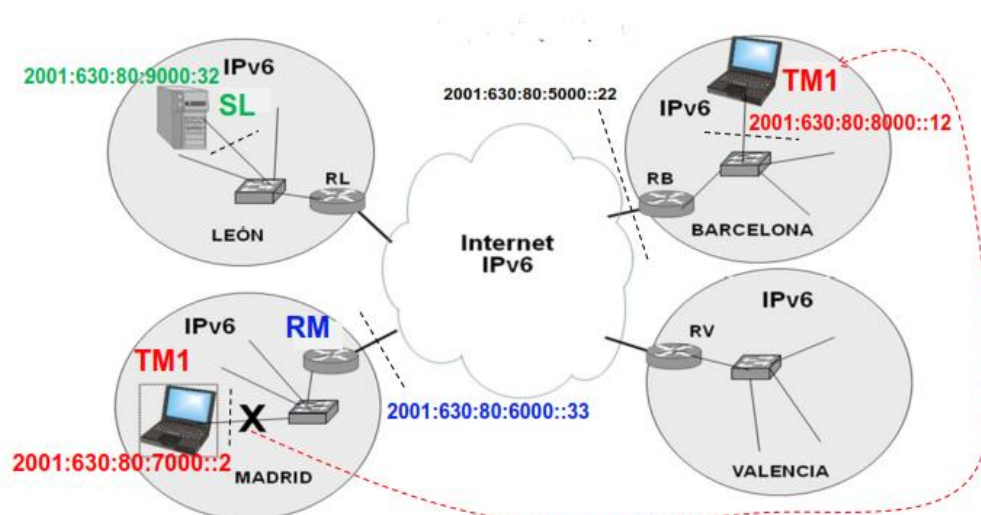
SOLUCIÓN.**Respuesta del inciso1.**

Figura 116: Solucion1 de MIPv6

- RM y TM1 deben tener funcionalidad de IP móvil, RB por el contrario no necesita ninguna funcionalidad adicional.
- RM debe registrar la asociación de la dirección permanente (2001:630:80:7000::2) y CoA (2001:630:80:8000::12) de TM1 a partir de una cabecera de extensión de movilidad recibida de TM1.
- RM debe encapsular los paquetes recibidos con dirección destino la dirección permanente de TM1 (2001:630:80:7000::2) en una nueva cabecera con dirección destino la CoA de TM1 (2001:630:80:8000::12) y dirección origen la dirección de RM del interfaz de Internet (2001:630:80:6000::33).
- TM1 debe ENVIAR a RM, para su registro, la asociación de la dirección permanente (2001:630:80:7000::2) y CoA (2001:630:80:8000::12) en una cabecera de extensión de movilidad.
- TM1, en caso de que el tráfico pase por RM, debe DESENCAPSULAR los paquetes recibidos con dirección destino la dirección CoA de TM1 (2001:630:80:8000::12).

- TM1, en caso de que el tráfico venga directamente de SL, debe CAMBIAR la dirección destino del paquete recibido (CoA: 2001:630:80:8000::12) por la dirección que está en la Cabecera de Extensión de encaminamiento (dirección permanente: 2001:630:80:7000::2).

Respuesta de inciso 2.

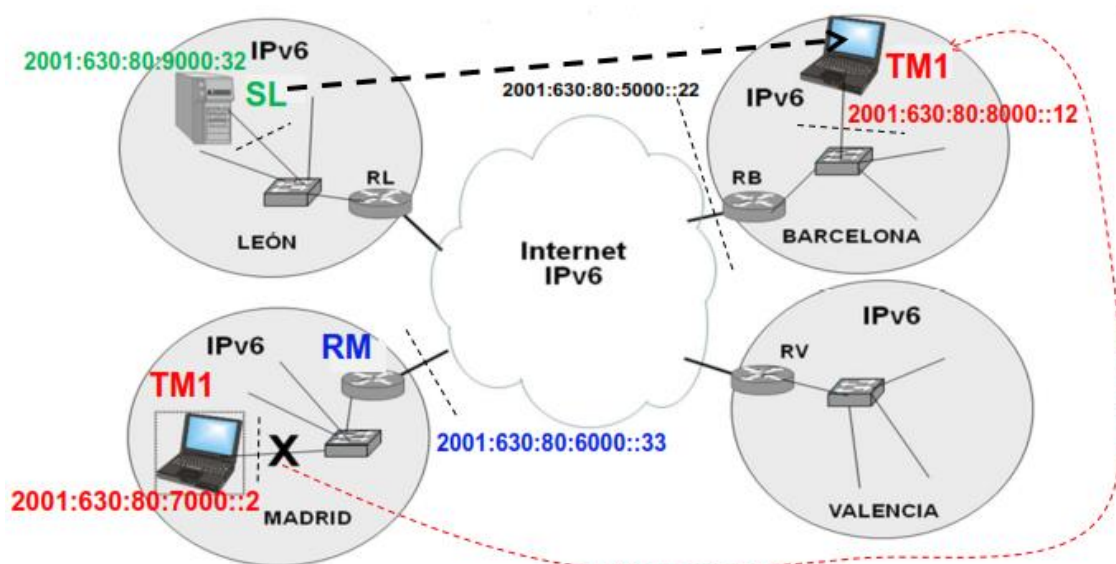


Figura 117: Solución del inciso 2 de MIPv6.

- Dirección Origen (SL): 2001:630:80:9000::32.
- Dirección Destino (CoA de TM1): 2001:630:80:8000::12.
- Cabecera siguiente: Cabecera de encaminamiento (43).
- Cabecera de Extensión de Encaminamiento: Dirección permanente de TM1 en la red de Madrid: 2001:630:80:7000::2.

Repuesta del inciso 3.

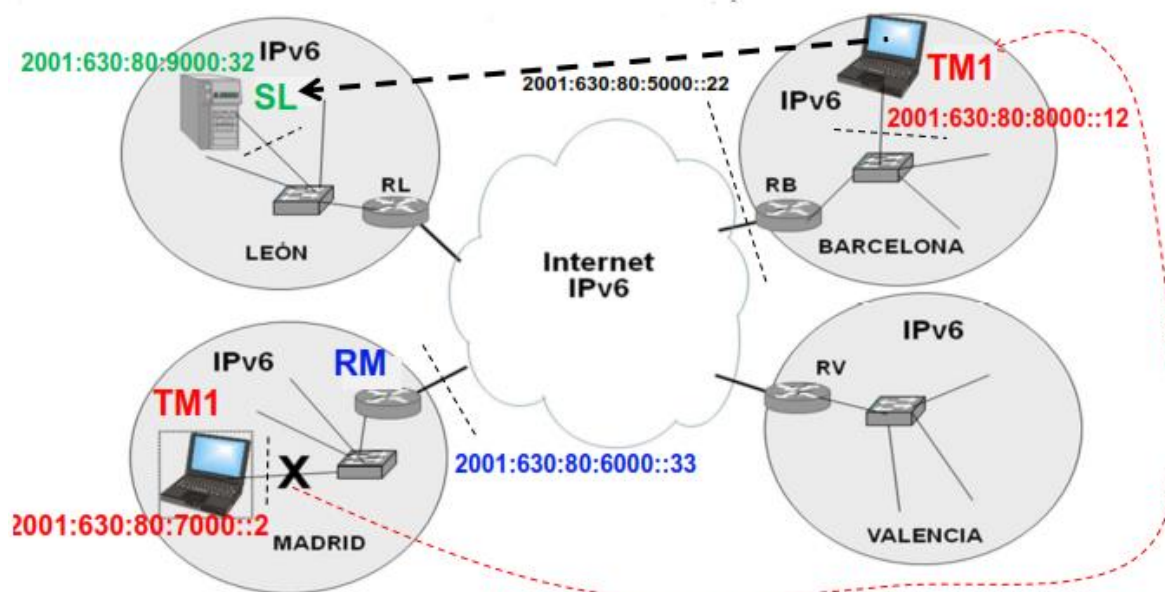


Figura 118: Solución del inciso3 MIPv6.

- Dirección Origen (CoA de TM1): 2001:630:80:8000::12.
- Dirección Destino (SL): 2001:630:80:9000::3.
- Cabecera siguiente: cabecera de extensión de opciones para el Destino (60).
- Cabecera de Extensión de Opciones para el Destino: Dirección permanente de TM1 en la red de Madrid: 2001:630:80:7000::2.

17. Seguridad en IPv6

17.1 IPsec

IPv6 incluye explícitamente la posibilidad de utilizar el modelo de seguridad IPsec (Internet Protocol Security) que **proporciona autenticidad, integridad y confidencialidad a las comunicaciones de extremo a extremo.**

IPsec es un conjunto de protocolos abiertos que tienen como fin proporcionar seguridad en las comunicaciones de la capa de red del modelo OSI (a la que pertenece el protocolo IPv6), y de ese modo, a todos los protocolos de capas superiores.

En IPv4 la implementación de IPsec se define en una especificación diferente a la del propio protocolo IPv4, por lo que la inclusión del protocolo se hace con mecanismos definidos fuera del mismo, mientras que en IPv6 la propia arquitectura "extensible" del protocolo permite implementar IPsec de forma natural. Es importante reseñar que IPv6 habilita la posibilidad de usar IPsec, y no los mecanismos de cifrado y autenticación propios de IPsec.

17.1.1 Modos de funcionamiento de IPsec.

IPsec tiene dos modos de funcionamiento que proporcionan distintos niveles de seguridad:

🚦 **Modo Transporte:** se cifra y/o autentica la carga útil, o payload, pero las cabeceras no se tienen en cuenta. Tiene como ventaja que se puede utilizar de extremo a extremo pero, por contra, la información de las cabeceras, como la dirección IP de origen y destino, es visible.

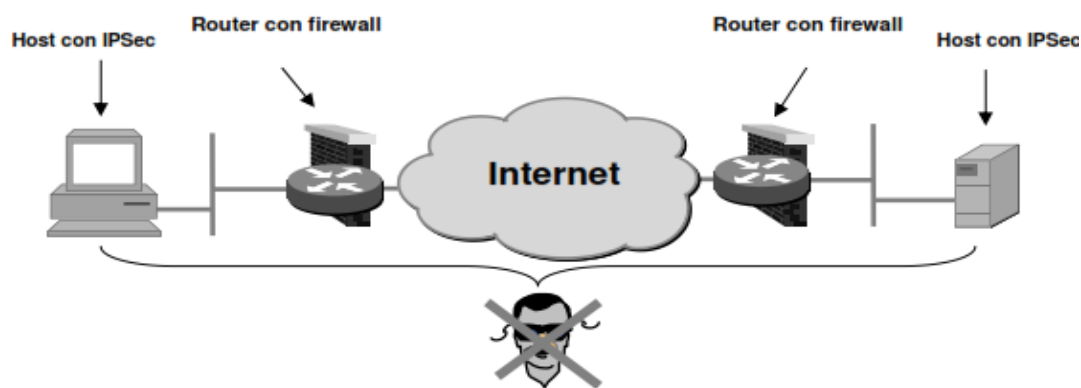


Figura 119: IPsec modo transporte.

En la figura se muestra un esquema en el cual se encuentran dos hosts de los cuales para que funcione el IPsec en modo transporte, se tiene que implementar IPsec en ambos hosts, el cual se tiene una comunicación segura de extremo a extremo, nos muestra dos routers con firewall de los cuales ya se ha implementado el IPsec desde los hosts y con esto se asegura de un extremo al otro se tenga seguridad.

✚ **Modo Túnel:** una plataforma, o pasarela, encapsula el paquete original en otro paquete. Con ello se cifra y/o autentica el paquete original completo, pero se necesita de una plataforma que realice el túnel.

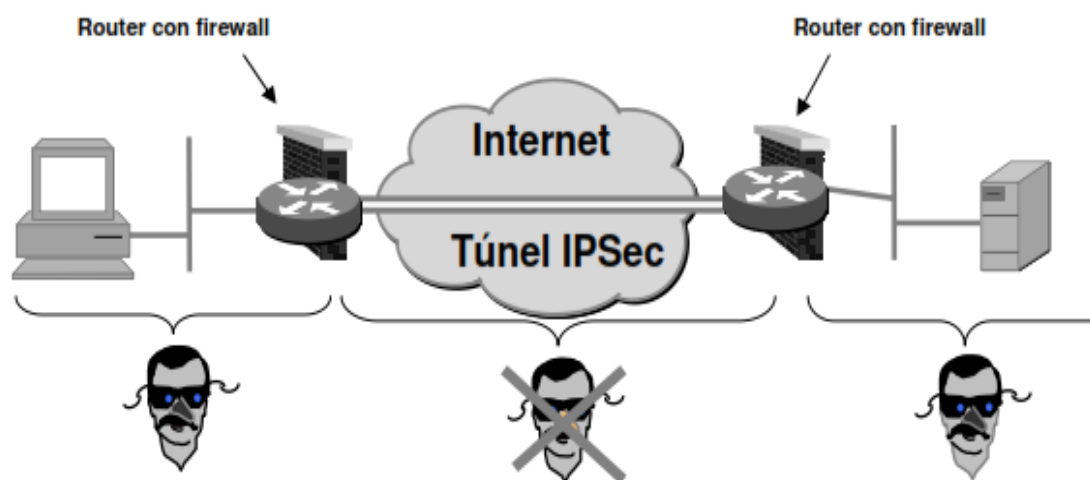


Figura 120: IPsec Modo Túnel.

En la figura se muestra un esquema en el cual se encuentran dos *hosts*, como también dos *routers* de los cuales para que funcione el IPsec en modo túnel, se tiene que implementar IPsec en ambos *routers* los cuales ejecutan una pasarela de seguridad. Este modo de funcionamiento de IPsec permite incorporarlo sin tener que modificar los *hosts*.

Además, IPsec tiene dos modos o protocolos de transferencia, que a su vez pueden funcionar en modo túnel o transporte:



➤ **AH (Authentication Header):** proporciona autenticación, integridad y un servicio de anti-repetición opcional.

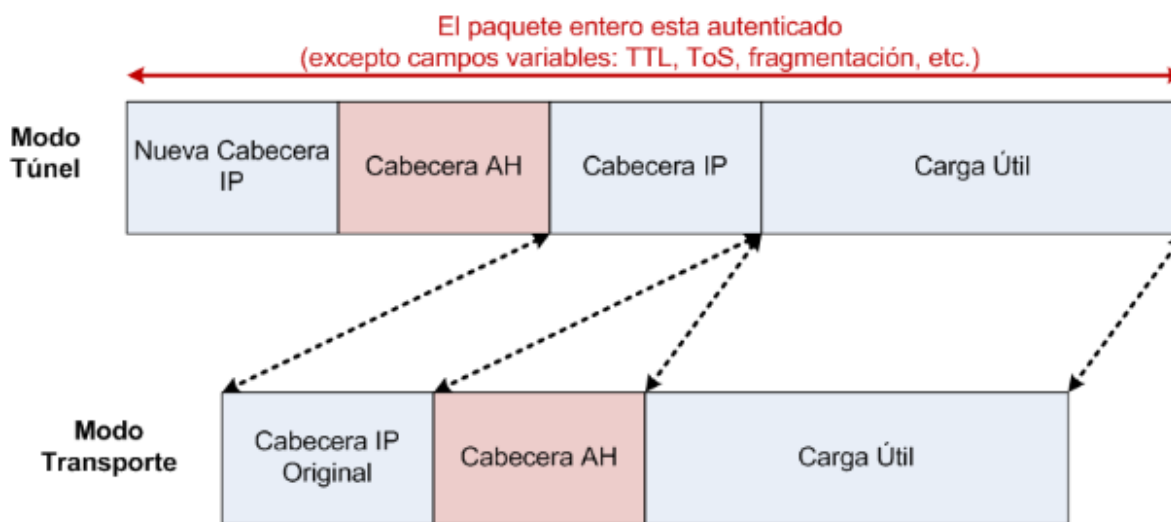


Figura 121: Implementación AH en modo túnel y modo transporte.

✓ AH en Modo Túnel

En éste modo la cabecera interna posee el origen y destino finales, mientras que la cabecera externa posee direcciones distintas (las de las puertas de enlace).

La cabecera AH protege a toda la cabecera interna, incluida la totalidad de la cabecera IP interna. La posición de la AH respecto a la cabecera IP externa es la misma que en el modo transferencia.

✓ AH en Modo Transporte

AH se inserta después de la cabecera IP, y antes del protocolo de capa superior (TCP, UDP, ICMP, ETC.) o antes de cualquier cabecera propia de IPSec que ya se haya incluido.

➤ **ESP (Encapsulating Security Payload):** además de las ventajas anteriores proporciona confidencialidad.

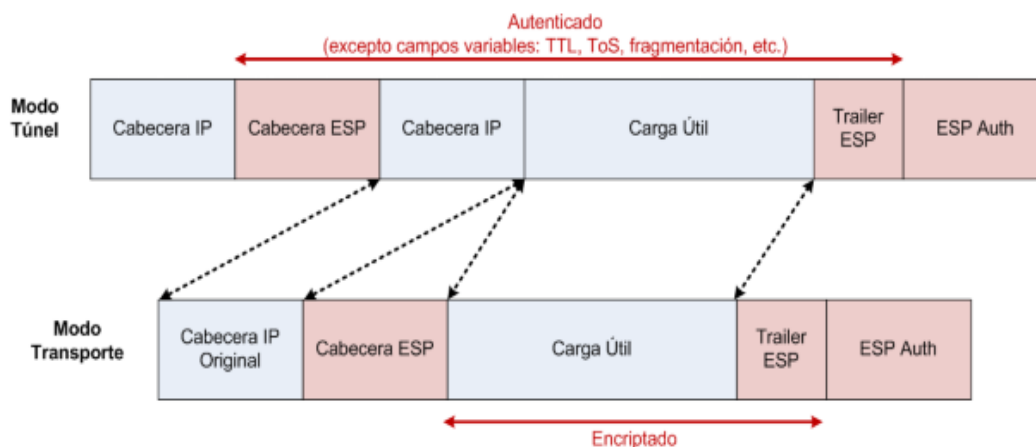


Figura 122: Implementación ESP en modo túnel y modo transporte.

✓ **ESP en Modo Túnel**

En éste modo la cabecera interna posee el origen y destino finales, mientras que la cabecera externa posee direcciones distintas (las de las puertas de enlace). ESP protege a toda la cabecera interna, incluida la totalidad de la Cabecera IP interna. La posición de la AH respecto a la cabecera IP externa es la misma que en el modo transferencia.

✓ **ESP en Modo Transporte**

ESP se inserta después de la cabecera IP, y antes del protocolo de capa superior (TCP, UDP, ICMP,...) o antes de cualquier cabecera propia de IPsec que ya se haya incluido.

17.1.2 Mayor fortaleza en la Red

La nueva versión del protocolo introduce novedades que mejoran la eficiencia del proceso de enrutamiento de los paquetes IP. Lo que permitirá que los elementos de red puedan gestionar mayor número de transmisiones y con mayor rapidez. Los cambios son los siguientes:

- ✓ Cabeceras simplificadas y de tamaño fijo.
- ✓ No se realiza fragmentación de los paquetes por los elementos intermedios. El tamaño de los paquetes los deberán determinar los extremos de la comunicación. Sin embargo, aunque a largo plazo debería acabar favoreciendo al flujo de datos, esta característica, al ser tan diferente de lo que se hace en IPv4, y al asentarse sobre ICMP, está provocando problemas en la implementación de IPv6, dando lugar a



errores de conectividad que están creando la impresión de que IPv6 no funciona completamente bien en la práctica.

- ✓ Facilita la agregación en las tablas de enrutamiento debido al uso estricto de CIDR para todos los tipos de direcciones, y a la mejor organización de sus asignaciones. Por otra parte, esta mejora es imprescindible debido al gran aumento de direcciones IP que se produce.
- ✓ Implementación obligatoria y mejorada del direccionamiento multicast. También se ha creado el direccionamiento anycast, en el que un grupo de servidores que proporcionan un mismo servicio comporten la misma dirección, de tal forma que el servidor seleccionado para dar dicho servicio vendrá determinado por la eficiencia de acceso. Aunque este direccionamiento es difícil de implementar en la práctica y, en su mayoría, es utilizado únicamente por los enrutadores.
- ✓ Utilización de etiquetas para QoS en las comunicaciones: el protocolo incluye la posibilidad de etiquetar clases y flujos de comunicaciones para que los enrutadores prioricen unas transmisiones sobre otras.

Otras mejoras.

- Imposibilidad de exploración de redes mediante "fuerza bruta". Anteriormente, los atacantes o programas maliciosos, como los gusanos, podían encontrar objetivos en una red comprobando todas las direcciones posibles. Pero debido al crecimiento exponencial de su número total, esta exploración es, a priori, inviable.
- Desaparece la necesidad de utilizar NAT. Aunque ha sido una tecnología muy útil, tiene los inconvenientes de que genera una falsa sensación de seguridad y de que se pierde la posibilidad de realizar conexiones seguras de extremo a extremo, incrementando la complejidad y el coste del desarrollo de aplicaciones.
- Se elimina la posibilidad de realizar un ataque DDOS de tipo broadcast o smurf⁶ al desaparecer este direccionamiento y al implantarse medidas de seguridad en el multicast.

17.1.3 Consideraciones de Seguridad en IPv6

Por el momento, el número de problemas de seguridad y ataques sobre IPv6 es pequeño debido a que no está desplegado aún a gran escala. Pero, se espera que la tendencia cambie a medida que los operadores y proveedores de contenidos lo implementen en sus redes y servicios.



En el siguiente apartado se describen los principales aspectos relacionados con la seguridad del protocolo que hay que tener en cuenta desde tres puntos de vista:

- **Aspectos técnicos.**
- **Consideraciones de gestión.**
- **Estructura o características propias del protocolo.**

1. Aspectos Técnicos.

✓ **Dispositivos de seguridad que no analizan el protocolo IPv6.**

Puede que los dispositivos de seguridad, como cortafuegos o IDS, o las herramientas de gestión de red no sean capaces o no estén configurados para analizar los flujos de datos del protocolo IPv6. Si fuera así, se podrían establecer comunicaciones maliciosas desde o hacia equipos de la red que soporten IPv6.

✓ **Presencia de dispositivos de los que se desconoce que pueden usar IPv6 y de túneles IPv6.**

Muchos Sistemas Operativos tienen habilitado IPv6 por defecto, como la mayor parte de sistemas Windows modernos, Mac OS X, Linux y Solaris.

Además pueden existir túneles IPv6. Un túnel es una conexión punto a punto, en la que se encapsulan los paquetes IPv6 en paquetes IPv4, de forma que se pueda transmitir IPv6 a través de una infraestructura IPv4. En el extremo final del túnel, se desencapsula (o extrae) el paquete IPv6 original.

Los dispositivos de seguridad perimetral puede que no estén preparados o configurados para analizar estos flujos de datos, que pueden ser utilizados para comunicaciones no permitidas como, por ejemplo, puertas traseras de C&C (Command and Control) de botnets o de P2P.

La posibilidad de crear túneles IPv6 se encuentra presente en todos los sistemas operativos, como Windows Vista y Windows 7 que tienen habilitado por defecto la tecnología Teredo, aunque se deshabilita si detecta que el equipo pertenece a un dominio o tiene soporte IPv6 a través de su red local. Otras formas de implementar túneles que pueden estar presentes son 6to4 e ISATAP.

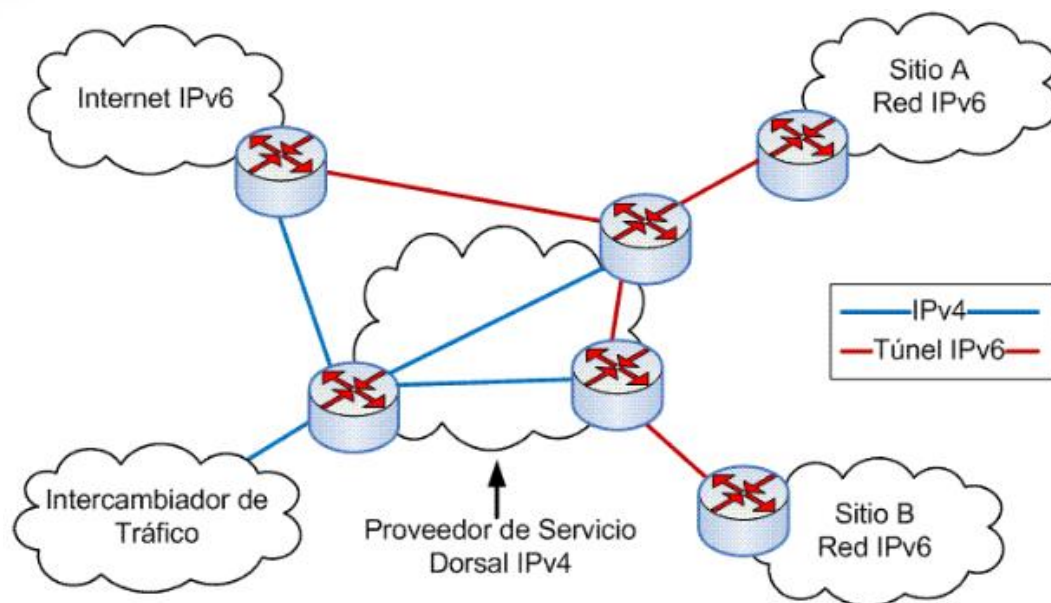


Figura 123: Túnel IPv6 en IPv4.

✓ Dejar de utilizar NAT

Una consecuencia indirecta del uso de NAT es que se emplea a modo de cortafuegos para proteger los equipos internos de las conexiones externas. Pero ya que IPv6 elimina su necesidad, se deberá modificar la política de los cortafuegos para que, según la política de seguridad, filtre o no las comunicaciones directas a los equipos de la red privada.

✓ Necesidad de multicast e ICMP

Muchos cortafuegos bloquean estos protocolos, aunque ciertas partes pueden ser muy importantes como, por ejemplo, el uso de ICMP para PMTU. En IPv6 son imprescindibles para su funcionamiento; por lo tanto, se deberán modificar las políticas de seguridad para permitir determinadas comunicaciones multicast e ICMP.

✓ Cambio en la monitorización de la red

Debido al gran número de direcciones disponibles será inviable escanear la red por fuerza bruta, por lo que los equipos se inventariarán de otra manera como, por ejemplo, en un servidor DNS.

Sin embargo, posiblemente surjan otras formas de escanear una red: existen direcciones multicast concretas para localizar servicios (por ejemplo FF05::2 All routers, FF05::1:3 All DHCP Servers) y direcciones de link-local, que permiten la comunicación en el segmento de red al que se esté conectado. Un atacante puede utilizar estas direcciones para establecer contacto con equipos o servicios. Aunque, en la práctica, este método no será probable que tenga éxito ya que la mayor parte de los sistemas operativos están configurados para no responder a estas peticiones.



✓ **Doble exposición IPv6 e IPv4.**

Durante años convivirán sistemas con doble pila, que soporten ambas versiones del protocolo, y mecanismos de transición a IPv6, lo que provocará que haya mayores posibilidades de existencia de vulnerabilidades.

Por otra parte, un sistema podrá ser atacado utilizando IPv4, IPv6 o una combinación de ambos, por ejemplo, utilizando IPv4 para detectar el equipo e IPv6 como canal oculto de comunicaciones.

✓ **Actualización de protocolos y equipos a IPv6.**

La gran mayoría de protocolos han sido adaptados para que utilicen direcciones IPv4 e IPv6, como por ejemplo BGP o DNS. La implantación de IPv6 supondrá la instalación y/o configuración de estos protocolos.

Existe el problema de que algunas aplicaciones que trabajan con IPv6 no son actualizadas frecuentemente. También existe actualmente una falta de soporte por parte de algunos fabricantes de enrutadores, switches y cortafuegos, aunque se prevé un mayor impulso a medida que haya una mayor adopción del protocolo.

2. Consideraciones de gestión.

✓ **Curva de Aprendizaje.**

Como con toda adopción de una nueva tecnología, las organizaciones necesitan de tiempo y recursos a la hora de adquirir el conocimiento necesario para implantar y administrar con seguridad el protocolo IPv6.

✓ **Implementación de sistemas de doble pila**

La implantación de IPv6 supondrá un importante cambio en los sistemas de comunicaciones ya que deberá soportar ambos protocolos y su interoperabilidad. El diseño, implantación y configuración de estos sistemas de doble pila, que implementan IPv4 e IPv6, será un proyecto complejo en el que habrá que evaluar todos los requisitos posibles de seguridad.

3. Estructura o características propias del protocolo.

El uso de IPv4 ha evolucionado con el tiempo, solucionándose los problemas que han surgido debido a su uso generalizado durante muchos años. De este modo se han creado tecnologías como NAT, CIDR o IPsec.

IPv6 puede que sufra un proceso similar, aunque atenuado por la experiencia que se posee con IPv4. Un ejemplo de este proceso de evolución del protocolo, es la decisión de que se rechacen los paquetes que utilizan la cabecera RH0, utilizada para definir la ruta de los paquetes, porque se podía utilizar para realizar un ataque de denegación de servicio.

Para los puntos descritos a continuación ya hay soluciones disponibles, aunque falta que algunos sistemas operativos las implementen.



✓ **Suplantación de identidad en la autoconfiguración de la dirección IP.**

Una de las novedades del protocolo es la capacidad de una interfaz de generar su dirección IP a partir de su dirección MAC. Durante este proceso el dispositivo pregunta al resto de dispositivos de la red si alguno está utilizando esa dirección. Además, si el dispositivo está conectado a una red en la que existe un enrutador, que dará el resto de parámetros de configuración como puede ser el prefijo de la red.

Durante este proceso, cualquier dispositivo podría generar de forma continuada una respuesta falsa informando de que la dirección está en uso y provocar que el dispositivo que solicita una dirección no se pueda conectar a la red. También, podría hacerse pasar por un enrutador para realizar un ataque de man-in-the-middle.

El protocolo SEND soluciona este problema, aunque todavía no ha sido implementado en la mayoría de sistemas operativos. SEND es una extensión que mejora la seguridad de protocolo NDP, que es el encargado de descubrir otros nodos en la red local, enrutadores, etc. Para realizar sus funciones SEND utiliza encriptación asimétrica y firma electrónica.

SEND es una evidente mejora respecto a IPv4 donde no existe nada comparable.

✓ **Privacidad**

Al generar un equipo su dirección IP a partir de la dirección MAC, se puede asociar una IP a un equipo de forma unívoca y, a su vez, se puede asociar un equipo a una persona.

Al realizar uso de Internet se deja un rastro de la dirección IP en los distintos servidores o redes con lo que se establece una comunicación. A partir de esta dirección IP se podría saber que servidores web o servicios visitó una persona.

Una solución para este problema consiste en la generación aleatoria de parte de la dirección IP, lo que se conoce como extensiones de privacidad. La gran mayoría de los sistemas operativos soportan las extensiones de privacidad y en algunos incluso están habilitadas por defecto (Windows XP, Vista y 7). Otra posible solución es la asignación temporal de direcciones mediante DHCPv6.

17.2 VPN

Una VPN o Red Privada Virtual, es una tecnología que permite la extensión de una red pública como Internet a un espacio de red local. Las VPN funcionan de manera tal que, si bien se utiliza una red pública como es la de conexión a Internet, los datos son transmitidos por un canal privado, de forma que no peligra la seguridad ni la integridad de la información interna. Los datos son cifrados y descifrados alternativamente.

17.2.1 Conexiones VPN a través de Internet IPv4

Para la mayor parte de las intranets actuales, las conexiones VPN se crean a través de Internet IPv4. En la figura 110 se muestran los componentes basados en Windows para conexiones VPN de este tipo. Estos componentes constan de lo siguiente:

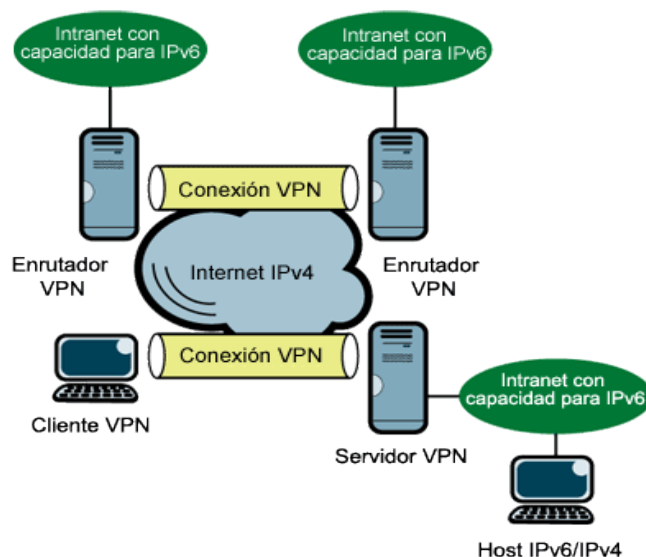


Figura 124: Componentes para conexiones VPN a través de Internet IPv4

- **Cliente VPN:** Se trata de un equipo que inicia una conexión VPN de acceso remoto a un servidor VPN y se comunica con recursos de intranet. Una conexión VPN de acceso remoto permite al cliente VPN actuar como si estuviera conectado directamente a la intranet. Un cliente VPN puede ejecutar versiones de cliente o servidor de Windows.
- **Servidor VPN:** Este equipo está a la escucha de intentos de conexión VPN remotos, aplica requisitos de autenticación y conexión y dirige paquetes entre clientes VPN y recursos de intranet. Un servidor VPN ejecuta normalmente una versión de servidor de Windows con el servicio de enrutamiento y acceso remoto.
- **Enrutador VPN:** Un enrutador VPN es un equipo que inicia o escucha intentos de conexión VPN de sitio a sitio. Una conexión VPN de sitio a sitio interconecta dos partes de una Intranet. Un enrutador VPN ejecuta una versión de servidor de Windows y el servicio de enrutamiento y acceso remoto.
- **Conexión VPN:** Una conexión VPN es el vínculo lógico entre el cliente VPN y el servidor VPN o bien, entre enrutadores VPN tal como se define en la encapsulación de un protocolo VPN.



- **Intranet habilitada con IPv6:** Esta intranet puede reenviar tráfico IPv6, de manera nativa o cómo túnel en forma de paquetes IPv4.
- **Host IPv6/IPv4:** Este nodo de intranet envía y recibe tráfico IPv6, nativamente o como túnel en forma de paquetes IPv4.

NOTA: Para conexiones VPN a través de Internet IPv4, hay dos métodos que se usan para enviar IPv6: paquetes IPv6 enviados por túnel en forma de paquetes IPv4, a partir de ahora denominados tráfico IPv6 sobre IPv4, y tráfico IPv6 nativo.

17.2.2 Tráfico IPv6 sobre IPv4

En este método, un cliente de acceso remoto o un host IPv6/IPv4 en la intranet encapsula paquetes IPv6 con un encabezado IPv4 y envían el resultado en forma de un paquete IPv4. Para intranets, la tecnología de transición IPv6 (RFC 4214) de ISATAP permite nodos IPv6/IPv4 para intercambiar tráfico IPv6 a través de una intranet de sólo IPv4. Con ISATAP, puede habilitar la conectividad IPv6 en la intranet sólo de IPv4 sin tener que configurar o actualizar los enrutadores existentes para admitir la asignación de direcciones y el reenvío de IPv6 nativo.

Para el tráfico IPv6 sobre IPv4, la carga del paquete IPv4 enviado a través de la conexión VPN es un paquete IPv6. La figura 111 muestra la estructura de paquete general para el tráfico VPN al enviar un paquete IPv6 sobre IPv4 usando una conexión VPN a través de Internet IPv4.

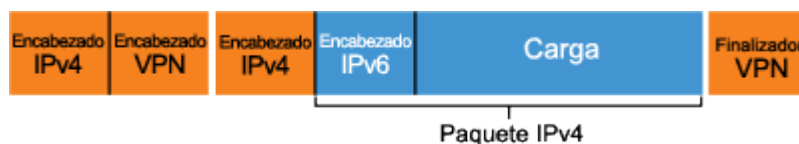


Figura 125: Paquetes IPv6 sobre IPv4 que usan una conexión VPN a través de Internet IPv4

17.2.3 Tráfico IPv6 nativo

Para el tráfico IPv6 nativo, el cliente VPN, el servidor o el enrutador envían paquetes IPv6 a través de la conexión VPN sin la encapsulación IPv4 inicial. Esto funciona para intranets con conectividad IPv6 nativa y requiere que los clientes VPN, servidores y enrutadores admitan el protocolo de control IPv6 (IPv6CP), RFC 2472, que define cómo los nodos IPv6 negocian opciones de configuración IPv6 para las conexiones basadas en el protocolo punto a punto (PPP).

La figura 112 muestra la estructura de paquete general para el tráfico VPN al enviar un paquete IPv6 usando una conexión VPN a través de Internet IPv4.



Figura 126: Paquetes IPv6 nativo que usan una conexión VPN a través de Internet IPv4

17.2.4 Conexiones VPN a través de Internet IPv6

También puede realizar conexiones VPN a través de Internet IPv6. Dichas conexiones VPN son raras ahora pero pasarán a ser más predominantes conforme más proveedores de servicios de Internet ofrezcan IPv6 a sus clientes y más organizaciones incluyan la conectividad a Internet IPv6 en sus redes perimetrales de intranet.

Para admitir conexiones VPN a través de Internet IPv6, los protocolos VPN que se usan deben admitir conexiones sobre IPv6. Las conexiones VPN a través de Internet IPv6 usan el mismo conjunto de componentes que los de las conexiones VPN a través de Internet IPv4 para las conexiones de acceso remoto y VPN de sitio a sitio.

También hay dos maneras de enviar paquetes IPv6 sobre Internet IPv6: Tráfico IPv6 sobre IPv4 y tráfico IPv6 nativo. La figura 113 muestra la estructura general de paquetes IPv6 sobre IPv4 cuando se envían a través de una conexión VPN a través de Internet IPv6.



Figura 127: Paquetes IPv6 sobre IPv4 que usan una conexión VPN a través de Internet IPv6

De la misma manera que para el tráfico IPv6 sobre IPv4 sobre Internet IPv4, el tráfico IPv6 sobre IPv4 a través de Internet IPv6 requiere la implementación de una tecnología de transición IPv6 como ISATAP en su intranet.

La figura muestra la estructura general de paquetes IPv6 cuando se envían a través de una conexión VPN a través de Internet IPv6. Al igual que para el tráfico IPv6 nativo sobre Internet IPv4, el tráfico IPv6 nativo sobre Internet IPv6 requiere compatibilidad con IPv6CP y la implementación de la conectividad IPv6 nativa en la intranet.



Figura 128: Paquetes IPv6 nativo que usan una conexión VPN a través de Internet IPv6

17.3 VPN sobre IPsec

Una red privada virtual (Virtual Private Network) es uno de los más implementados, donde mediante un proceso de encapsulación, y en este caso, de encriptación, de los paquetes de datos a distintos puntos remotos mediante el uso de infraestructuras públicas de transporte.

En el caso de acceso remoto, la VPN permite al usuario acceder a su red corporativa, asignándole a su ordenador remoto las direcciones y privilegios de la misma, aunque la conexión la haya realizado por medio de un acceso a Internet público, todo esto con la seguridad que provee IPsec para la transportación de sus datos.

Las implementaciones de VPN conjuntamente con IPsec representan una gran solución para las empresas en cuanto a seguridad, confidencialidad e integridad de los datos.

Reduce significativamente el costo de la transferencia de datos de un lugar a otro, con una buena implementación y configuración de IPsec las empresas sacan provecho de todas las ventajas que el IPsec provee.



Figura 129: Método de trabajo de IPsec en IPv6, redes privadas virtuales

Road Warrior

Los trabajadores que pasan gran parte del tiempo fuera de la empresa, como teletrabajadores o los llamados “road warriors”, personas que necesitan acceder a la red de la empresa pero que están constantemente cambiando de ubicación.

Ahora lo que se permite es a un ordenador personal o portátil el acceso a la red corporativa, manteniendo la privacidad.

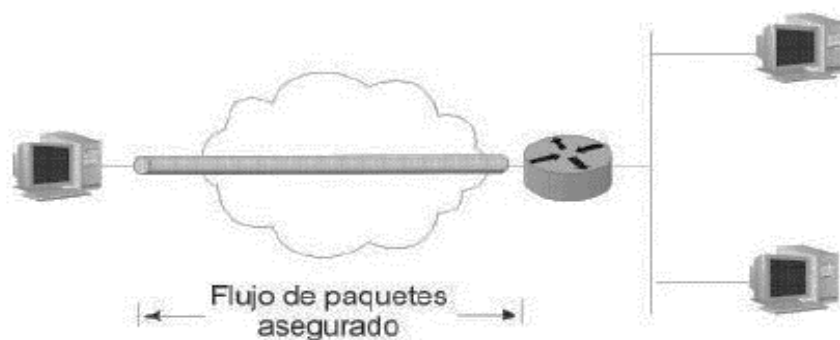


Figura 130: Método de trabajo de IPsec en IPv6, road warrior

Túneles anidados

Este método de trabajo de IPsec, no es muy recomendable, por lo complicado de su construcción, mantenimiento y consume de recursos de red.

Un ejemplo que se puede describir en la figura 117 es, el host A envía un paquete al host B, la política indica que debe ser autenticado con el router RB, donde existe una VPN entre el router RA y RB.

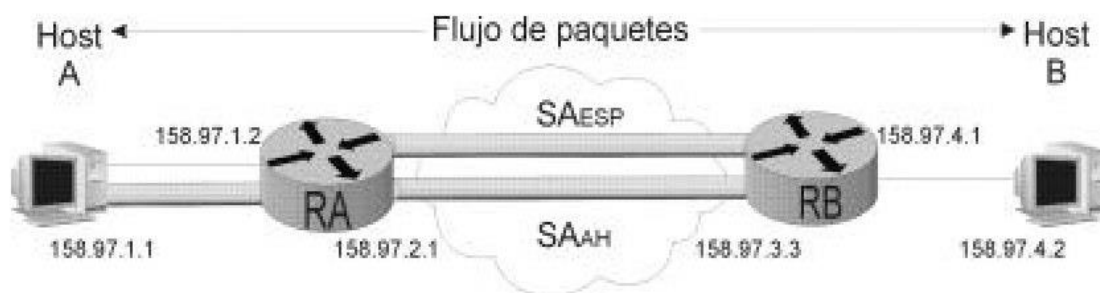


Figura 131: Método de trabajo de IPsec en IPV6, túneles anidados



18. Servicios en IPv6

18.1 DNS

Domain Name System o DNS es un sistema de nomenclatura jerárquica para computadoras, servicios o cualquier recurso conectado a Internet o a una red privada. Este sistema asocia información variada con nombres de dominios asignados a cada uno de los participantes. Su función más importante, es traducir (resolver) nombres inteligibles para las personas en identificadores binarios asociados con los equipos conectados a la red, esto con el propósito de poder localizar y direccionar estos equipos.

El DNS como base de datos:

El objetivo principal del DNS es entonces almacenar información de mapeo entre nombres y números IP

- Directa y reversa

El sistema opera entonces como una base de datos distribuida en la que existe la posibilidad de delegar su administración de sectores del espacio de nombres a diferentes organizaciones.

Componentes

Para la operación práctica del sistema DNS se utilizan tres componentes principales:

- Los **Clientes fase 1**: Un programa cliente DNS que se ejecuta en la computadora del usuario y que genera peticiones DNS de resolución de nombres a un servidor DNS (Por ejemplo: ¿Qué dirección IP corresponde a nombre. Dominio?).
- Los **Servidores DNS**: Que contestan las peticiones de los clientes. Los servidores recursivos tienen la capacidad de reenviar la petición a otro servidor si no disponen de la dirección solicitada.
- Y las **Zonas de autoridad**, es una parte del espacio de nombre de dominios sobre la que es responsable un servidor DNS, que puede tener autoridad sobre varias zonas. (Por ejemplo: subdominio .ORG, .COM, etc).

Servicios Proporcionados por DNS

Existen dos formas de identificación de un host:

- Por un nombre de host.
- Por una dirección IP.
- Las personas prefieren la identificación de nombre de host, mientras que los routers prefieren las direcciones IP, de longitud fija y estructurada jerárquicamente.
- Una base de datos distribuida implementada en una jerarquía de servidores de nombres.
- Una aplicación de la capa de aplicación que permite que se comuniquen los host y los servidores de nombres para proporcionar el servicio de traducción.



- DNS es empleado comúnmente por otros protocolos de la capa de aplicación (incluyendo HTTP, SMTP,FTP) para traducir los nombres de host proporcionados por los usuarios a direcciones IP.
- La propia máquina del usuario ejecuta el lado cliente de la aplicación DNS.
- El navegador extrae el nombre de host su URL y lo pasa al lado cliente de la aplicación DNS.
- El cliente DNS envía una consulta a un servidor DNS con el nombre del host, lo que constituye el mensaje DNS de petición.
- El cliente DNS eventualmente recibe una respuesta que incluye la dirección IP para el nombre del host.

Alias de host.

Un nombre de host con un nombre complejo puede tener uno o más nombres de alias. Por ejemplo, el nombre de host clear.comp.ejemplo.com podría tener dos alias, como ejemplo.com y www.ejemplo.com. En este caso, se dice que host clear.comp.ejemplo.com se dice que es el nombre canónico. Los alias de nombres de host, de existir, son más nemotécnicos que el nombre canónico. El DNS puede ser invocado por una aplicación (proporcionando un nombre alias) el nombre canónico del host y su dirección IP.

Alias de servidor de correo.

Por razones obvias, es recomendable que las direcciones de correo electrónico sean mnemotécnicas. El DNS puede ser invocado por una aplicación de correo para obtener el nombre canónico de host a partir del alias proporcionado, así como la dirección IP del host.

Distribución de carga

El DNS es también utilizado para realizar una distribución de carga entre servidores replicados, como los son los servidores web replicados. Cuando un cliente envía una consulta DNS para un nombre que tiene asociado un conjunto de direcciones, el servidor responde con el conjunto completo de direcciones IP, pero cambia el orden de las direcciones en cada respuesta. Dado que el cliente típicamente envía un mensaje HTTP de petición a la primera dirección de la lista, la rotación DNS distribuye el tráfico sobre todos los servidores replicados. La rotación DNS también es utilizada en el correo electrónico, de forma que múltiples servidores de correo puedan tener el mismo nombre de alias.

Tipos de servidores DNS.

Servidor Primario/Principal/Maestro. Mantiene de forma oficial la información de la zona sobre la que tiene autoridad. Responde a las peticiones de resolución consultando sus propios archivos de zona.

Servidor Secundario/Esclavo. Obtiene la información de la zona pidiéndosela constantemente al servidor primario.

Servidor Cache. Se utilizan para acelerar las consultas de resolución de nombres de dominio frecuentemente utilizados. Suelen emplearse en redes de área local.



NOTA: El objetivo de tener varios servidores principales es distribuir la carga y dar cierta tolerancia a fallos. Cuando uno de los servidores principales falla, todas las peticiones acabarán en los demás. Por supuesto, este esquema no nos protege de fallos del servidor que produzcan errores en todas las peticiones DNS, como podrían ser errores del software.

Registros DNS.

Los servidores de nombres que conjuntamente implementan la base de datos distribuida DNS, almacenan **registros de recursos (RR)** para las correspondencias **nombre de host/dirección IP**, cada uno de estos mensajes DNS transportan uno o más registros de recursos.

Los registros de recursos son 4, y contienen los siguientes campos: (**Nombre, Tipo, Valor, TTL**).

El campo TTL es el tiempo de vida del registro de recursos, determina el momento en el que el recurso debe de ser borrado de la cache. En los ejemplos de registros siguientes ignoramos el campo TTL. El significado del nombre valor dependen del tipo:

- **AAAA** = Address – (dirección) Este registro se usa en IPv6 para traducir nombres de hosts a direcciones IPv6.
- **CNAME** = Canonical Name – (nombre canónico) Se usa para crear nombres de servidores de alojamiento adicionales, o alias, para los servidores de alojamiento de un dominio. Es usado cuando se están corriendo múltiples servicios (como FTP y servidor web) en un servidor con una sola dirección IP. Cada servicio tiene su propia entrada de DNS (como `ftp.ejemplo.com` y `www.ejemplo.com`). Esto también es usado cuando corres múltiples servidores HTTP, con diferentes nombres, sobre el mismo host. Se escribe primero el alias y luego el nombre real. Ej. `Ejemplo1 IN CNAME ejemplo2`
- **NS** = Name Server – (Servidor de Nombres) Define la asociación que existe entre un nombre de dominio y los servidores de nombres que almacenan la información de dicho dominio. Cada dominio se puede asociar a una cantidad cualquiera de servidores de nombres.
- **MX** = Mail Exchange – (registro de intercambio de correo) Asocia un nombre de dominio a una lista de servidores de intercambio de correo para ese dominio. Tiene un balanceo de carga y prioridad para el uso de uno o más servicios de correo.
- **PTR** = Pointer – (indicador) También conocido como 'registro inverso', funciona a la inversa del registro A, traduciendo IPs en nombres de dominio. Se usa en el archivo de configuración del DNS reversiva.
- **SOA** = Start of authority – (Autoridad de la zona) Proporciona información sobre el servidor DNS primario de la zona.



- **HINFO** = Host INFOrmation – (información del sistema informático) Descripción del host, permite que la gente conozca el tipo de máquina y sistema operativo al que corresponde un dominio.
- **TXT** = Text - (Información textual) Permite a los dominios identificarse de modos arbitrarios.
- **LOC** = Localización - Permite indicar las coordenadas del dominio.
- **WKS** - Generalización del registro MX para indicar los servicios que ofrece el dominio. Obsoleto en favor de SRV.
- **SRV** = Servicios - Permite indicar los servicios que ofrece el dominio. RFC 2782. Excepto MX y NS. Hay que incorporar el nombre del servicio, protocolo, dominio completo, prioridad del servicio, peso, puerto y el equipo completo.
- **SPF** = Sender Policy Framework - Ayuda a combatir el spam. En este registro se especifica cual o cuales hosts están autorizados a enviar correo desde el dominio dado. El servidor que recibe, consulta el SPF para comparar la IP desde la cual le llega con los datos de este registro.
- **ANY** = Toda la información de todos los tipos que exista.

Mensajes DNS.

El protocolo DNS define dos tipos de mensajes: query y reply, ambos con el mismo formato.

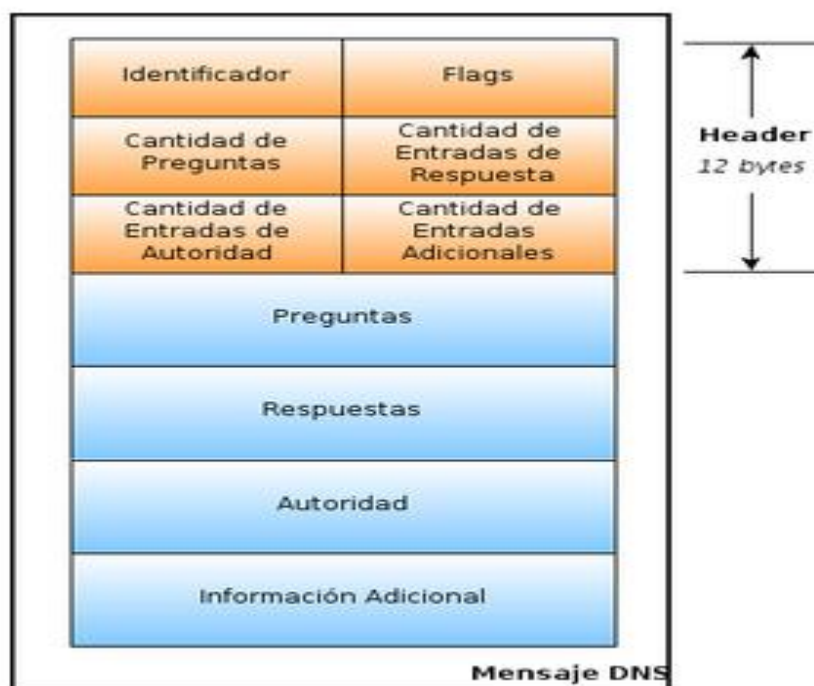


Figura 132: Mensaje DNS.



Representación gráfica de un mensaje DNS

Las secciones dentro del mensaje se organizan de la siguiente forma **Header**.

Los primeros 12 bytes forman la **sección de cabecera** (header section). A su vez, esta se compone de varios campos:

- Los primeros 16-bits son un número **identificador** de la query. Este identificador se copia al mensaje de reply de una query, de esta forma el cliente DNS puede juntar los mensajes de respuesta recibidos con las solicitudes enviadas.
- Luego viene un campo que tiene un conjunto de **flags**. El primer bit es una flag indicando si el mensaje es una query (0) o una reply (1).
- El segundo bit se coloca como 1 cuando el servidor DNS es autoritativo para el nombre solicitado.
- El tercer bit se marca como 1 cuando el cliente (que puede ser un host o un servidor DNS) desea que el servidor DNS utilice recursión en el caso de no tener la entrada pedida.
- El cuarto bit se coloca en 1 en una respuesta si el servidor DNS soporta recursión.
- Finalmente, hay cuatro campos que contienen el número de preguntas, respuestas, autoridad e información adicional (las siguientes secciones).

Questions

La sección de **preguntas (questions)** contiene información sobre la query. Contiene un campo del nombre que se está solicitando, y un campo de tipo.

Answers.

En una respuesta desde un servidor DNS, la sección de **respuestas (answers)** contiene las entradas asociadas al nombre que se solicitó originalmente. Puede tener múltiples entradas, debido a que un hostname puede tener múltiples IPs asociadas.

Authority.

La sección de **autoridad (authority)** contiene entradas de otros servidores autoritativos.

Additional Information.

La sección de **información adicional (additional information)** contiene otras entradas útiles que complementan a las respuestas anteriores.

Métodos de resolución de nombre.

Petición iterada.



Se dice que una consulta o petición es iterada cuando todas las respuestas de la petición son devueltas directamente al mismo servidor que las envió.

Para entender pensemos en un ejemplo donde un host desea conocer una dirección. Tenemos un servidor DNS local `dns.usm.cl`, y el host desea conocer la dirección de `lugares.wikipedia.com`, en donde un servidor DNS autoritativo para `lugares.wikipedia.com` es `dns.wikipedia.com`. Cuando el host envía el un mensaje de consulta DNS a su servidor DNS local, en este caso `dns.usm.cl`, el mensaje de consulta tiene el nombre de host que debe ser traducido, es decir `lugares.wikipedia.com`. El servidor DNS local reenvía la consulta a un servidor raíz, el cual toma el sufijo `.com`, y devuelve al servidor DNS local una lista de las direcciones IP de los servidores TLD responsables del dominio `.com`. El servidor DNS local reenvía nuevamente el mensaje de consulta a uno de estos servidores TLD, en donde el servidor TLD revisa el sufijo `wikipedia.com` y responde la dirección IP del servidor autoritativo, es decir `dns.wikipedia.com`. Por último, el servidor DNS local reenvía la consulta directamente a `dns.wikipedia.com`, que responde con la dirección IP de `lugares.wikipedia.com`. Luego el servidor DNS local cuenta con la dirección IP correspondiente. Como se puede observar se realizan peticiones iteradas, en donde cada petición la realiza el servidor local DNS, al cual se le devuelven cada resultado. Se puede observar en detalle en la imagen "Consultas iterativas en DNS".

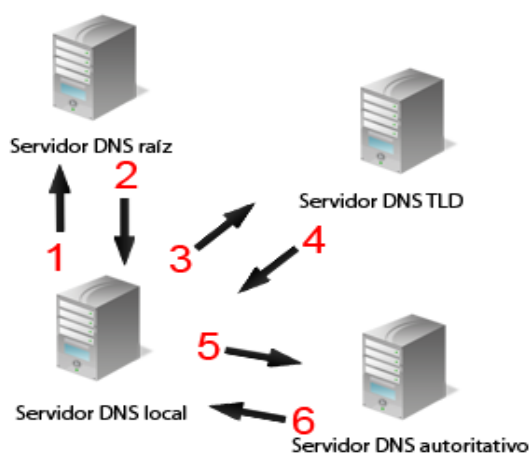


Figura 133: Ejemplo de petición iterada.

Petición recursiva.

Se realiza una petición o consulta recursiva, cuando se solicita a un servidor que obtenga por sí mismo la correspondencia.

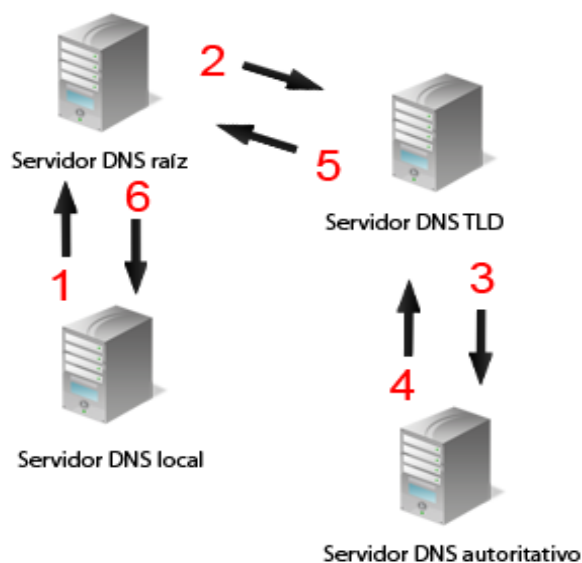


Figura 134: Ejemplo de petición recursiva.

Consultas recursivas en DNS.

Para entender mejor el concepto, se aplica el siguiente ejemplo. Un host alumno.usm.cl desea conocer una dirección. Tenemos un Servidor DNS local dns.usm.cl, y el host desea conocer la dirección de lugares.wikipedia.com, en donde un servidor DNS autoritativo para lugares.wikipedia.com es dns.wikipedia.com. El host envía un mensaje de consulta DNS a su servidor DNS local, dns.usm.cl, en donde el mensaje de consulta contiene el nombre de host que debe ser traducido, es decir lugares.wikipedia.com. El servidor DNS local consulta al servidor raíz, a su vez este al servidor DNS TLD, el cual consulta al servidor DNS autoritativo, dns.wikipedia.com para obtener la dirección IP de lugares.wikipedia.com. Obtenida la información se retorna al servidor DNS TLD, este a su vez al servidor raíz, y finalmente al servidor DNS local dns.usm.cl él envía al host solicitante. Se observan 8 procedimientos los cuales se pueden observar en la figura "Consultas recursivas DNS".

Caché.

Como realizar la llamada a un servidor de DNS lejano puede agregar un delay considerable, se utiliza caché, guardando la información de las entradas DNS en servidores locales, los que se sincronizan periódicamente con los Root DNS Servers.

18.2 FTP

Servidor Very Secure FTP en IPv6.

Basado en UNIX, VSFTPD (Very Secure FTP Daemon) es usado para implementar servidores de archivos a través del protocolo FTP. Se diferencia porque su configuración por defecto es muy segura. En la actualidad, VSFTPD es considerado uno de los servidores FTP más seguros del mundo. Trabaja sobre TCP/IP e incluye soporte para IPv6. Utiliza los puertos 20 y 21 en TCP.



Servicios de administración remota para FTP en IPv6

Los servicios de administración remota son importantes, sobre todo para la configuración remota de routers y otros dispositivos. Los servicios SSH y telnet, son compatibles con Windows, IOS, Linux y la mayoría de los sistemas Unix.

18.3 HTTP en IPv6.

Hypertext Transfer Protocol o HTTP (en español protocolo de transferencia de hipertexto) es el protocolo usado en cada transacción de la World Wide Web. HTTP fue desarrollado por el World Wide Web Consortium y la Internet Engineering Task Force.

Es un protocolo orientado a transacciones y sigue el esquema petición-respuesta entre un cliente y un servidor. Al cliente que efectúa la petición (un navegador web o un spider) se lo conoce como "user agent" (agente del usuario). A la información transmitida se la llama recurso y se la identifica mediante un localizador uniforme de recursos (URL). Los recursos pueden ser archivos, el resultado de la ejecución de un programa, una consulta a una base de datos, la traducción automática de un documento, etc.

HTTP es un protocolo sin estado, es decir, que no guarda ninguna información sobre conexiones anteriores.

La representación textual definida para las direcciones IPv6, no es directamente compatible con Uniform Resource Locator URL ya que usa ":" y "." como caracteres delimitadores. El RFC 2396 propone que para utilizar una dirección IPv6 en una URL, la dirección literal, debe ser encerrada entre corchete.

Por ejemplo las direcciones IPv6:

FEDC: BA98: 7654:3210: FEDC: BA98: 7654:3210

1080:0:0:0:8:800:200C: 4171

3FFE: 2A00: 100:7031 :: 1

Se representaría con el siguiente ejemplo:

http:// [FEDC: BA98: 7654:3210:FEDC:BA98:7654:3210]:80/index.html

http:// [1080:0:0:0:8:800:200C:417A] / index.html

http:// [3FFE: 2A00: 100:7031 :: 1]

Transacciones HTTP

Una transacción HTTP está formada por un encabezado seguido, opcionalmente, por una línea en blanco y algún dato. El encabezado especificará cosas como la acción requerida del servidor, o el tipo de dato retornado, o el código de estado.



El servidor puede excluir cualquier encabezado que ya esté procesado, como Authorization, Content-type y Content-length. El servidor puede elegir excluir alguno o todos los encabezados, si incluirlos, si se excede algún límite del entorno de sistema. Ejemplos de esto son las variables HTTP_ACCEPT y HTTP_USER_AGENT.

- **HTTP_ACCEPT.** Los tipos MIME que el cliente aceptará, dados los encabezados HTTP. Otros protocolos quizás necesiten obtener esta información de otro lugar. Los elementos de esta lista deben estar separados por una coma, como se dice en la especificación HTTP: tipo, tipo.
- **HTTP_USER_AGENT.** El navegador que utiliza el cliente para realizar la petición. El formato general para esta variable es: software/versión biblioteca/versión.

Servidor Apache (HTTP) en IPv6.

Apache, es un servidor web Hypertext Transfer Protocol HTTP de código abierto, es compatible con plataformas Unix (Berkeley Software Distribution BSD, GNU/Linux, etc.), Macintosh, Microsoft Windows entre otras. Usado esencialmente para enviar páginas web estáticas y dinámicas. Dentro de sus características importantes, el soporte para IPv6 en sus configuraciones, lo cual permite la adaptabilidad a nuevas funciones y características. Trabaja sobre el puerto 80 en Transmission Control Protocol TCP.

Los cambios necesarios cuando se utiliza Apache solo con IPv6 o para IPv4/IPv6 son de acuerdo a la necesidad del tipo de servidor a implementarse. Si, en la configuración se especifica que el servidor debe funcionar con una dirección IPv4 en particular, se necesita actualizar la configuración para incluir una dirección IPv6. Dentro del archivo /etc/httpd/conf/httpd.conf, la directiva Listen 80, al momento de activar IPv6 requiere el valor de [::]:80, tener en cuenta que la dirección IPv6 está encerrado entre corchetes. Con esto el servidor escucha todas las direcciones IPv4 e IPv6, a menos que IPv4 este desactivada, se escuchará sólo direcciones IPv6. Si se desea activar IPv4, se debe agregar la directiva Listen 0.0.0.0:80, la línea habilita escuchar todo el pool de direcciones IPv4.

Servidor Very Secure FTP en IPv6.

Basado en UNIX, VSFTPD (Very Secure FTP Daemon) es usado para implementar servidores de archivos a través del protocolo FTP. Se diferencia porque su configuración por defecto es muy segura. En la actualidad, VSFTPD es considerado uno de los servidores FTP más seguros del mundo. Trabaja sobre TCP/IP e incluye soporte para IPv6. Utiliza los puertos 20 y 21 en TCP.

Servicios de administración remota para FTP en IPv6.

Los servicios de administración remota son importantes, sobre todo para la configuración remota de routers y otros dispositivos. Los servicios SSH y telnet, son compatibles con Windows, IOS, Linux y la mayoría de los sistemas Unix.



18.3 Secure Shell.

El protocolo SSH (Secure Shell) es una herramienta que nos permite conectarnos a equipos remotos (Servidores en Producción) así mismo, nos da la capacidad de llevar a cabo tareas administrativas dentro del mismo como, activar o apagar servicios. Además de la conexión a otros equipos, SSH nos permite copiar datos de forma segura, gestionar claves RSA para no escribir claves al conectar a las máquinas y pasar los datos de cualquier otra aplicación por un canal seguro tunelizado mediante SSH. Una clave RSA (Sistema Criptográfico con Clave Pública) es un algoritmo que genera un par de llaves de autenticación, la pública y la privada. La pública se distribuye en forma autenticada y la privada que generalmente es guardada en secreto por el propietario.

A diferencia de otros protocolos de comunicación remota tales como FTP o Telnet, SSH encripta la sesión de conexión, haciendo imposible que alguien pueda obtener contraseñas no encriptadas. De manera predeterminada, el protocolo SSH atiende peticiones por el puerto 22.

El protocolo SSH (Secure Shell) está implementado bajo el estándar TCP/IP, el cual a su vez se encuentra dividido en 5 secciones:

- Nivel Físico.
- Nivel De Enlace.
- Nivel de Internet.
- Nivel de Transporte.
- Nivel de Aplicación.

La capa de aplicación es el nivel que los programas más comunes utilizan para comunicarse a través de una red con otros programas. Los procesos que acontecen en este nivel son aplicaciones específicas que pasan los datos al nivel de aplicación en el formato que internamente use el programa y es codificado de acuerdo con un protocolo estándar.

SSH está diseñado para reemplazar los métodos más viejos y menos seguros para registrarse remotamente en otro sistema a través de la shell de comando, tales como **telnet** o **rsh**. Un programa relacionado, el **scp**, reemplaza otros programas diseñados para copiar archivos entre hosts como **rcp**. Ya que estas aplicaciones antiguas no encriptan contraseñas entre el cliente y el servidor, evite usarlas mientras le sea posible. El uso de métodos seguros para registrarse remotamente a otros sistemas reduce los riesgos de seguridad tanto para el sistema cliente como para el sistema remoto.

Cómo funciona el protocolo SSH



Figura 135: Funcionamiento del protocolo SSH.

El funcionamiento de este protocolo se puede resumir en los siguientes pasos que dejamos a continuación:

1. **El cliente inicia una conexión TCP sobre el puerto 22 del servicio.** Este puerto es el que utiliza por defecto el protocolo, aunque como veremos en siguientes puntos, se puede modificar.
2. **El cliente y el servidor se ponen de acuerdo en la versión del protocolo a utilizar**, así como el algoritmo de cifrado utilizado para el intercambio de la información.
3. El **servidor**, que tiene en su poder dos claves (una privada y una pública), **manda su clave pública al cliente**.
4. Cuando el cliente recibe la clave enviada por el servidor, la compara con la que tiene almacenada para verificar su autenticidad. **El protocolo SSH exige que el cliente la confirme la primera vez.**
5. Con la clave pública del servidor en su poder, **el cliente genera una clave de sesión aleatoria**, creando un mensaje que contiene esa clave y el algoritmo seleccionado para la encriptación de la información. Toda esa información es enviada al servidor haciendo uso de la clave pública que envió en un paso anterior de forma cifrada.
6. Si todo es correcto, el cliente queda autenticado, iniciando la sesión para comunicarse con el servidor.



Características.

El protocolo SSH proporciona los siguientes tipos de protección:

- Después de la conexión inicial, el cliente puede verificar que se está conectando al mismo servidor al que se conectó anteriormente.
- El cliente transmite su información de autenticación al servidor usando una encriptación robusta de 128 bits.
- Todos los datos enviados y recibidos durante la sesión se transfieren por medio de encriptación de 128 bits, lo cual los hacen extremadamente difícil de descifrar y leer.
- El cliente tiene la posibilidad de reenviar aplicaciones X11 desde el servidor. Esta técnica, llamada *reenvío por X11*, proporciona un medio seguro para usar aplicaciones gráficas sobre una red.

Ya que el protocolo SSH encripta todo lo que envía y recibe, se puede usar para asegurar protocolos inseguros. El servidor SSH puede convertirse en un conducto para convertir en seguros los protocolos inseguros mediante el uso de una técnica llamada *reenvío por puerto*, como por ejemplo POP, incrementando la seguridad del sistema en general y de los datos.

Importancia SSH.

Los usuarios tienen a su disposición una variedad de herramientas que les permiten interceptar y redirigir el tráfico de la red para ganar acceso al sistema. En términos generales, estas amenazas se pueden catalogar del siguiente modo:

- Intercepción de la comunicación entre dos sistemas — En este escenario, existe un tercero en algún lugar de la red entre entidades en comunicación que hace una copia de la información que pasa entre ellas. La parte interceptora puede interceptar y conservar la información, o puede modificar la información y luego enviarla al recipiente al cual estaba destinada.

Este ataque se puede montar a través del uso de un paquete sniffer — una utilidad de red muy común.

- Personificación de un determinado host — Con esta estrategia, un sistema interceptor finge ser el recipiente a quien está destinado un mensaje. Si funciona la estrategia, el sistema del usuario no se da cuenta del engaño y continúa la comunicación con el host incorrecto.

Esto se produce con técnicas como el envenenamiento del DNS o spoofing de IP (engaño de direcciones IP).

Si se utiliza SSH para inicios de sesión de shell remota y para copiar archivos, se pueden disminuir estas amenazas a la seguridad notablemente. Esto es porque el cliente SSH y el servidor usan firmas digitales



para verificar su identidad. Adicionalmente, toda la comunicación entre los sistemas cliente y servidor es encriptada. No servirán de nada los intentos de falsificar la identidad de cualquiera de los dos lados de la comunicación ya que cada paquete está cifrado por medio de una llave conocida sólo por el sistema local y el remoto.

18.4 WiFi (802.11).

El estándar 802.11 establece los niveles inferiores del modelo OSI para las conexiones inalámbricas que utilizan ondas electromagnéticas, por ejemplo:

- La capa física (a veces abreviada capa "PHY") ofrece tres tipos de codificación de información.
- La capa de enlace de datos compuesta por dos subcapas: **control de enlace lógico (LLC)** y **control de acceso al medio (MAC)**.

La capa física define la modulación de las ondas de radio y las características de señalización para la transmisión de datos mientras que la capa de enlace de datos define la interfaz entre el bus del equipo y la capa física, en particular un método de acceso parecido al utilizado en el estándar Ethernet, y las reglas para la comunicación entre las estaciones de la red.

El estándar 802.11 tiene tres capas físicas que establecen modos de transmisión alternativos:

Capa de enlace de datos (MAC)	802.2
	802.11
Capa física	DSSS FHSS Infrarrojo

Figura 136: capas del estándar 802.11

Cualquier protocolo de nivel superior puede utilizarse en una red inalámbrica Wi-Fi de la misma manera que puede utilizarse en una red Ethernet.

Los distintos estándares 802.11

El estándar 802.11 en realidad es el primer estándar y permite un ancho de banda de 1 a 2 Mbps. El estándar original se ha modificado para optimizar el ancho de banda (incluidos los estándares 802.11a, 802.11b y 802.11g, denominados estándares físicos 802.11) o para especificar componentes de mejor manera con el fin de garantizar mayor seguridad o compatibilidad. La tabla a continuación muestra las distintas modificaciones del estándar 802.11 y sus significados:

Los estándares 802.11a, 802.11b y 802.11g, llamados "estándares físicos", son modificaciones del estándar 802.11 y operan de modos diferentes, lo que les permite alcanzar distintas velocidades en la transferencia de datos según sus rangos.



ESTANDAR	CARACTERISTICAS
802.11a	• Cuenta con velocidad de transmisión de 2Mbps. • Opera en la banda 5Ghz y utiliza 52 subportadoras con una velocidad máxima de 54Mbps/s lo que lo hace un estándar practica para redes inalámbricas con velocidad reales de aproximadamente 20Mbps/s. • Tiene 12 canales sin solapa, 8 para red inalámbrica y 4 para conexiones Access point.
802.11b	• Tiene una velocidad máxima de 11Mbps/s. • Utiliza el mismo método de acceso definido en el estándar original CSMA/cs. • Funciona en la banda de 2.4Ghz. • La velocidad máxima de conexión con este estándar es aproximadamente 5.9Mbps/s sobre TCP y 7.1Mbps/s sobre UDP.
802.11g	• Utiliza la banda de 2.4Ghz. • Opera a una velocidad teórica máxima de 54Mbps/s que en promedio es de 22Mbps/s de velocidad real de transferencia. • Es compatible con el estándar B y utiliza las mismas frecuencias.
802.11n	• Puede trabajar en dos dispositivos de frecuencia 2.4Hz y 5Ghz. • Compatible con dispositivos basados en todas las ediciones anteriores de WI-FI, permite alcanzar un mayor rendimiento en una velocidad de 600Mbps en capa física

Figura 137: Distintos estándares 802.11.

Tipos de Seguridad.

La seguridad es el punto débil de las redes inalámbricas, pues la señal se propaga por el aire en todas las direcciones y puede ser captada a distancia de centenares de metros, utilizando una notebook con antena. Esto hace que las redes inalámbricas sean vulnerables a ser interceptadas.

WEP.

Es el algoritmo opcional de seguridad para brindar protección a las redes inalámbricas, incluido en la primera versión del estándar IEEE 802.11, mantenido sin cambios en las nuevas 802.11a y 802.11b, con el fin de garantizar compatibilidad entre distintos fabricantes. El WEP es un sistema de encriptación estándar implementado en la MAC y soportado por la mayoría de las soluciones inalámbricas. En ningún caso es compatible con IPSec.

Cifrado.

WEP utiliza una clave secreta compartida entre una estación inalámbrica y un punto de acceso. Todos los datos enviados y recibidos entre la estación y el punto de acceso pueden ser cifrados utilizando esta clave compartida. El estándar 802.11 no especifica cómo se establece la clave secreta, pero permite que haya



una tabla que asocie una clave exclusiva con cada estación. En la práctica general, sin embargo, una misma clave es compartida entre todas las estaciones y puntos de acceso de un sistema dado.

Autenticación:

WEP proporciona dos tipos de autenticación:

- un sistema abierto, en el que todos los usuarios tienen permiso para acceder a la WLAN
- Autenticación mediante clave compartida, que controla el acceso a la WLAN y evita accesos no autorizados a la red.

Características:

- Proporciona confidencialidad, autenticación y control de acceso en redes WLAN
- utiliza una misma clave simétrica y estática en las estaciones y el punto de acceso.
- El estándar no contempla ningún mecanismo de distribución automática de claves, lo que obliga a escribir la clave manualmente en cada uno de los elementos de red.

Algoritmo RC4.

Es un algoritmo de Cifrado de flujo (no de bloques), creado en 1987 por Ronald Rivest (la R de RSA - Secreto Comercial de RSA Data Security). Funciona a partir de una clave de 1 a 256 bytes (8 a 1024 bits), inicializando una tabla de estados. Esta tabla se usa para generar una lista de bytes pseudo-aleatorios, los cuales se combinan mediante la función XOR con el texto en claro; el resultado es el texto cifrado.

Fallas de seguridad.

Debilidad del vector de inicialización

La implementación del vector de inicialización (IV) en el algoritmo WEP tiene varios problemas de seguridad. Recordemos que el IV es la parte que varía de la clave (*seed*) para impedir que un posible atacante recopile suficiente información cifrada con una misma clave.

Sin embargo, el estándar 802.11 no especifica cómo manejar el IV; se indica que debería cambiarse en cada trama para mejorar la privacidad, pero no obliga a ello.

WPA (Wi-Fi Protected Access).

Surgió para corregir las limitaciones del WEP. Introdujo mejoras de seguridad tales como:

- TKIP (Temporal Key Integrity Protocol), que varía por sí solo la contraseña Wi-Fi cada cierto tiempo.
- Su variante más normal es la WPA-Personal. Usa el sistema PSK, o de clave precompartida. En él, todos los usuarios de la red inalámbrica tienen una misma contraseña Wi-Fi, que el propio usuario define. Ve más abajo cómo elegir una clave fuerte.



- También hay una versión WPA empresarial (*WPA-Enterprise*). Ofrece seguridad adicional al obligar al usuario a identificarse con un nombre y contraseña en sistemas de autenticación especiales, como RADIUS o 802.1X.

WPA2.

Es el estándar más moderno para proteger redes inalámbricas y el que recomienda la Wi-Fi Alliance. Existe también una versión personal (*WPA2-Personal*) y empresarial (*WPA2-Enterprise*).

WPA2 es compatible con WPA, lo que significa que en tu red Wi-Fi puedes usar PCs o dispositivos (router, adaptadores de red...) que admitan uno u otro sistema.

WPA2 no es compatible, sin embargo, con sistemas WEP. No podrás juntar en una misma red Wi-Fi dispositivos que sólo admitan WEP con otros válidos para WPA2. Es por razones de seguridad.

19. VOIP siguiente generación de voz en IPv6

VoIP o Voz sobre Protocolo de Internet (IP) es un sistema de telefonía que proporciona voz las llamadas telefónicas a través de redes de datos IP. La principal característica de esta tecnología basada en IP es que se envía como datos de conversaciones (o IP) a través de Internet.

En la actualidad, se está jugando un papel vital en la sustitución de la infraestructura de telefonía (basado en TDM) de hoy. Esta telefonía avanzada trae beneficios tanto a los consumidores como las empresas (o comerciales) de los clientes. La principal razón para migrar a VOIP es reducir el coste de la comunicación de voz (residencial y comercial).

19.1 Migración de voz IP versión 6

Para realizar la migración de las redes convencionales de voz a las redes IPv6, se debe describir la tecnología de señalización de voz sobre IP, donde la más óptima es la que usa el Protocolo de Inicio de Sesiones (SIP) definido en el RFC 3261 de la IETF.

Este protocolo SIP, se concentra en el establecimiento, modificación y terminación de las sesiones, y se complementa, entre otros, con el Protocolo de Descripción de Sesión (SDP), que describe el contenido multimedia de la sesión, por ejemplo qué direcciones IP, puertos y códecs se usarán durante la comunicación. También se relaciona directamente con el Real-time Transport Protocol (RTP), que es el verdadero portador del contenido de voz y vídeo que intercambian los participantes en una sesión establecida por SIP.

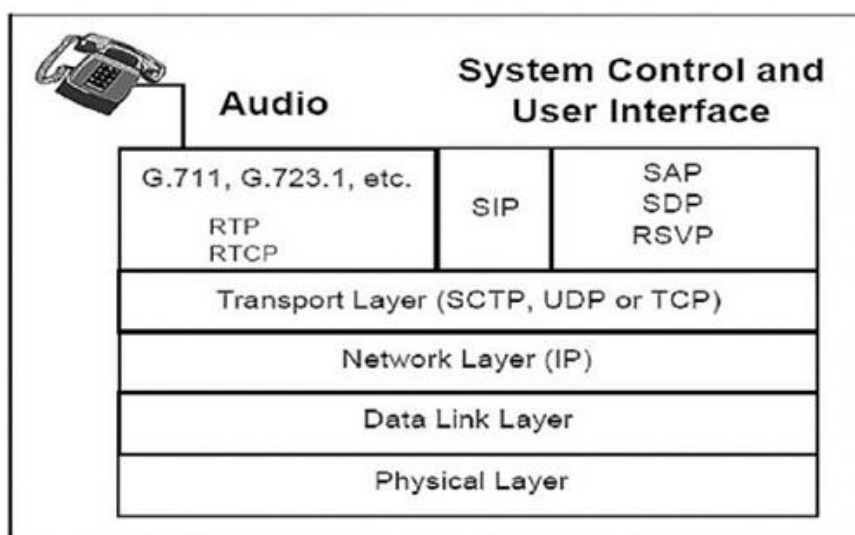


Figura 138: Pila de protocolo SIP

El SIP fue diseñado de acuerdo con el modelo de Internet por el grupo de trabajo Multiparty Multimedia Session Control (MMUSIC) del Internet Engineering Task Force (IETF), con el fin de estandarizar sesiones interactivas de usuario, con elementos multimedia; es un protocolo de señalización extremo a extremo que implica que toda la lógica se almacena en los dispositivos finales (salvo el enrutado de los mensajes SIP). El estado de la conexión también se almacena en los dispositivos finales.

El precio a pagar por esta capacidad de distribución y su gran escalabilidad es una sobrecarga en la cabecera de los mensajes, producto de tener que enviar toda la información entre los dispositivos finales.

Para que la transición hacia IPv6, mediante el protocolo SIP, sea una solución completa tiene que soportarse tanto en la capa de señalización como en la capa de sesión. Aunque SIP puede manejar redes heterogéneas IPv6/IPv4 en la capa de señalización, siempre que los servidores proxy y el Sistema de Nombres de Dominio (DNS) sean configurados correctamente, los agentes de usuario con diferentes redes y la dirección que usen diferentes redes y espacios de direcciones deben implementar las extensiones para el intercambio de voz entre ellos.

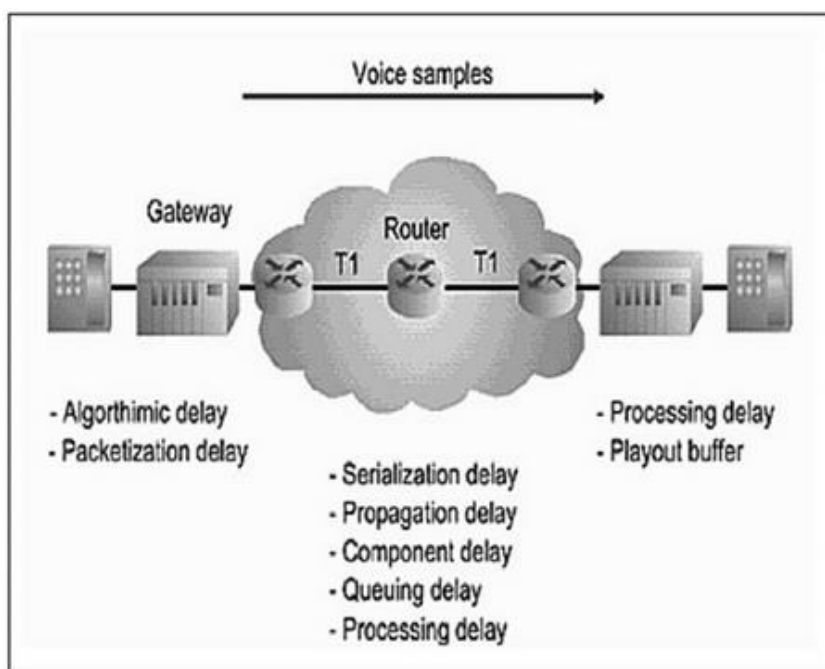


Figura 139: Fuentes de retardo de VoIP.



19.2 Calidad de servicio.

El principal problema que presentan las diferentes aplicaciones que funcionan sobre IP (entre ellas VoIP) es el relacionado con la garantía de QoS, ya que es necesario reservar anchos de banda y valores muy pequeños de fluctuación (jitter), retardo (delay), latencia, supresión de eco, etc. Particularmente, las diferentes fuentes de retardo en transmisión de VoIP se pueden observar en la figura 5, por eso, se presentan diversos problemas para garantizar la calidad del servicio.

Los problemas de la calidad del servicio en VoIP vienen derivados de dos factores principalmente: 1. Internet es un sistema basado en conmutación de paquetes y, por tanto, la información no viaja siempre por el mismo camino. Esto produce efectos como la pérdida de paquetes; 2. Las comunicaciones VoIP son en tiempo real, lo que hace que efectos como el eco, la pérdida de paquetes y latencia sean muy molestos y perjudiciales y deban ser evitados o controlados.

Cuando las tramas se transmiten a través de una red IP, la cantidad de retardo experimentado por cada una puede diferir. Esto lo causa la cantidad de retraso de encolamiento y tiempo de procesamiento, que puede variar dependiendo del tráfico cargado en la red. Sin embargo, como el gateway fuente genera tramas de voz a intervalos regulares (es decir, cada 20 ms), el Gateway destino típicamente no recibirá tramas de voz en intervalos regulares, lo que produce el jitter.

19.3 Desempeño IPv6 frente a IPv4

En redes IPv4, a fin de garantizar cierta "reserva", se hace uso del ítem "servicio diferenciado" dentro del encabezado. En IPv6, los campos Clase de Tráfico (que especifican parámetros de prioridad, retardo, rendimiento, fiabilidad) y Etiqueta de Flujo (que especifica la forma de etiquetar los paquetes de voz como pertenecientes a un mismo flujo), permiten al origen solicitar un manejo especial por parte de los nodos intermedios [10]. Si se configuran adecuadamente estos dos campos, se logra reducir el retardo de la transmisión y se genera una mayor calidad de audio, lo que repercute en el mejoramiento del servicio.

IPv6 proporciona mayor facilidad de clasificar los paquetes con identificadores de tráfico. Adicionalmente, el campo Etiqueta de Flujo tiene la ventaja de estar localizado antes de los campos de dirección, lo que ayuda a reducir los retardos en la verificación del paquete. En la Tabla 1 se muestran los beneficios que tiene IPv6 sobre IPv4:



Beneficio	IPv4	IPv6
Integridad punto a punto de la señalización de VoIP.		X
Seguridad (escucha disimulada)		X
Adaptabilidad	X	X
Fiabilidad		X
Alojamiento NAT (Network Address Translation)	X	
Calidad de servicio (QoS)		X
Soporte a tráfico multimedia en tiempo Real		X
Movilidad		X
Configuración dinámica		X

Tabla 30: Beneficios VoIPv6 respecto de VoIPv4

19.4 Arquitectura de integración.

Para lograr integración entre los servicios de voz proporcionados por las redes análogas, la Red Telefónica Pública Básica Conmutada (RTPBC), las redes bajo IPv4 y las redes que operan bajo IPv6, es necesario definir una arquitectura donde se puedan tener diferentes esquemas de direccionamiento y afrontar el cambio de versión del protocolo IP. Para solucionar esto se deben implementar puertas de enlace distribuidas, una pila doble que permita las dos versiones del protocolo y entidades de transición, además de un árbol con los pasos para la estrategia de transición. Una propuesta de arquitectura se presenta en la siguiente figura.

En la Figura se observa que hasta el Gateway la transferencia multimedia y la señalización se realizan a través de MTP, ISUP y DSS1, protocolos y elementos propios de la RTPBC. Cuando se hace la transición hacia redes IP, se hace a nivel de capas inferiores, en cualquiera de las versiones: 4 ó 6. En las capas superiores se utiliza SIP. En este punto hacen su aparición los nodos Media Gateway Control (MGC) y Media Gateway (MG): el primero se encarga de la adaptación de señalización, interfaz a la capa superior, gestión de políticas etc., y el segundo realiza la adaptación multimedia en los límites de la RTPBC y la red IP.

De aquí en adelante, se realiza el enrutamiento y la interoperabilidad de IP entre versiones 4 y 6 de manera normal; en el caso específico de la Figura, se emplea NAT como traductor.

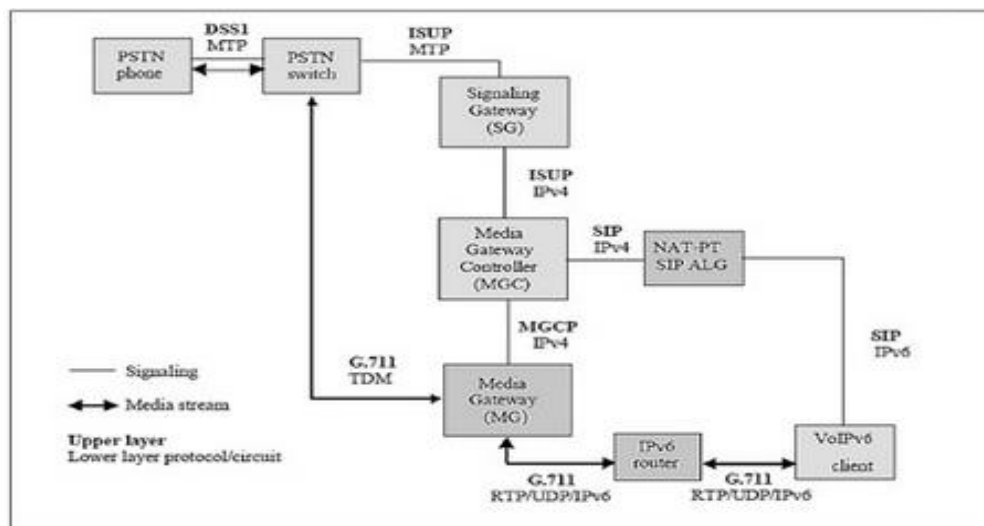


Figura 140: Arquitectura propuesta VoIPv6.

19.5 Medición del retardo

Para cuantificar la mejora de IPv6 respecto a IPv4, se realizaron llamadas utilizando el servicio Asteriskv6, que es un programa de software libre que implementa transmisión de voz sobre IPv6 y al cual se le puede configurar el protocolo SIP para que haga la transmisión sobre este. Se usó un servidor con Asteriskv6 y clientes Ubuntu con linphone y Windows con X-lite como software de voz.

Se tomaron como muestra diez llamadas y se realizaron diferentes capturas mediante Wireshark, donde se tomaron los paquetes recibidos y los tiempos de duración; los datos obtenidos se observan en la Tabla siguiente.

Paquetes VoIP transmitidos	Tiempo de envío (en segundos)	
	IPv4	Tiempo de envío (en segundos) IPv6
30123	465.6	237
40243	482.4	301.2
50132	656.4	407.4
60317	753	478.2
70322	948	593.4
80412	1084.2	604.8
90162	1315.2	781.8

100202	1350.0	883.8
110371	1488	952.2
120614	1626.6	1057.8

Tabla 31: Medición del retardo.

Estos tiempos se pueden graficar en forma de barras para observar la mejora sustancial en el tiempo del retardo, que se da cuando se emplea IPv6.

Para calcular el retardo se utiliza la ecuación:

$$\text{Retardo} = \frac{\text{tiempo de la llamada}}{\text{número de paquetes}}$$

Partiendo de los valores de la Tabla anterior y la Figura siguiente, se obtienen los retardos para los tiempos de la muestra:

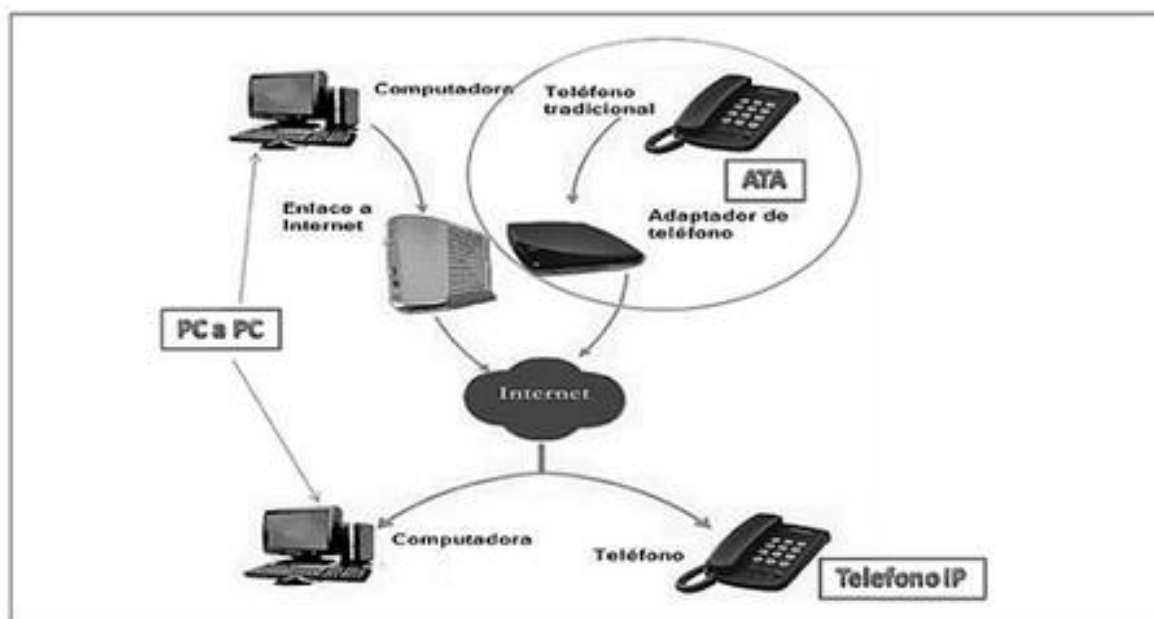


Figura 141: Tiempos de llamada vs. Paquetes IPv4 e IPv6.



CAPITULO III: DESARROLLO PRÁCTICO





ORGANIZACIÓN DE LAS PRÁCTICAS.

Introducción

La asignatura “Despliegue de IPv6” correspondiente a la Electiva X que pertenece al plan académico 2011 de la carrera de Ingeniería en Telemática del Departamento de Computación de la UNAN-León y se imparte en el segundo semestre del V año de la carrera e incluye en su currículo, 2 horas semanales de teoría y 2 horas semanales de laboratorio.

El contenido de este documento podrá ser revisado (de forma no significativa) para adaptar los trabajos y tareas del laboratorio a la evolución de la asignatura en su conjunto.

Programación

Práctica 1: Direccionamiento IPv6 con rutas estáticas.

Práctica 2: DHCPv6 y Autoconfiguración.

Práctica 3: Coexistencia de IPv4 e IPv6.

Práctica 4: VLANs estáticas y dinámicas.

Práctica 5: Frame Relay e Intervlans.

practica6: Listas de acceso.

Práctica 7: Enrutamiento interno: RIPng, OSPFv3, EIGRPv6.

Práctica 8: Enrutamiento dinámico: RIPng, OSPFv3, EIGRPv6, IS-ISv6 + BGP4

Práctica 9: VPN IPv6

Práctica 10: VoIP con IPv6.

Práctica 11: DNS.

Práctica 12: HTTP.

Práctica 13: SSH.

Practica 14: Miscelánea.

Evaluación

Los contenidos del laboratorio tendrán un valor del 40% en la calificación global de la asignatura. Los métodos de evaluación excluyentes entre sí:

Evaluación continua, será el método por defecto y se basará en:

- a. Control de las prácticas realizadas por parte del alumno en el laboratorio, que será específico para cada una de ellas o podrá consistir en preguntas individuales y concretas sobre su proceso de realización, ejecución y desarrollo de pequeños ejercicios de índole práctico, etc.



El resultado final de cada práctica será “apto” o “no apto”. Serán necesarios obtener al menos 4 “aptos” para poder presentarse al examen del Laboratorio.

Organización de las prácticas

Prácticas	Tiempo estimado de solución
Práctica 1: Direccionamiento IPv6 con rutas estáticas.	4 horas
Práctica 2: DHCPv6 y Autoconfiguración.	5 horas
Práctica 3: Coexistencia de IPv4 e IPv6.	7 horas
Práctica 4: VLANs estáticas y dinámicas.	5 horas
Práctica 5: Frame Relay e Intervlans.	6 horas
Práctica 6: Listas de acceso.	5 horas
Práctica 7: Enrutamiento interno: RIPng, OSPFv3, EIGRPv6.	8 horas
Práctica 8: Enrutamiento Dinámico: RIPng, OSPFv3, EIGRPv6 + IS-ISv6 + BGP4.	8 horas
Práctica 9: VPN IPv6	6 horas
Práctica 10: VoIP (Tunnel IPv4, IPv6)	10 horas
Práctica 11: DNS.	7 horas
Práctica 12: HTTP.	4 horas
Práctica 13: SSH.	4 horas
Práctica 14: Miscelánea.	14 horas
Total	93

NOTA: Las horas de tiempo estimado de solución incluye tantas horas presenciales y no presenciales. Esta distribución temporal es orientativa y podrá revisarse con objeto de permitir la realización de los trabajos previstos en todas las prácticas, así como las actividades de evaluación correspondientes.



PRÁCTICA 1: DIRECCIONAMIENTO IPv6 CON RUTAS ESTÁTICAS

Objetivo general

- Asignar direcciones IPv6 correctamente a los diferentes dispositivos terminales.

Objetivos específicos

- Agregar las rutas estáticas IPv6 de forma adecuadas a los routers para la comunicación entre diferentes redes LAN.
- Conocer cuáles son las direcciones IPv6 correctas de una dirección de Red para los equipos finales.

Introducción

En la siguiente práctica estudiaremos el funcionamiento de asignación de direcciones IPv6. En esta práctica se asignará direcciones IPv6 habilitables de una dirección de Red a los equipos finales pertenecientes a las Redes LAN, de igual manera se estudiara el enrutamiento estático IPv6

Requerimientos

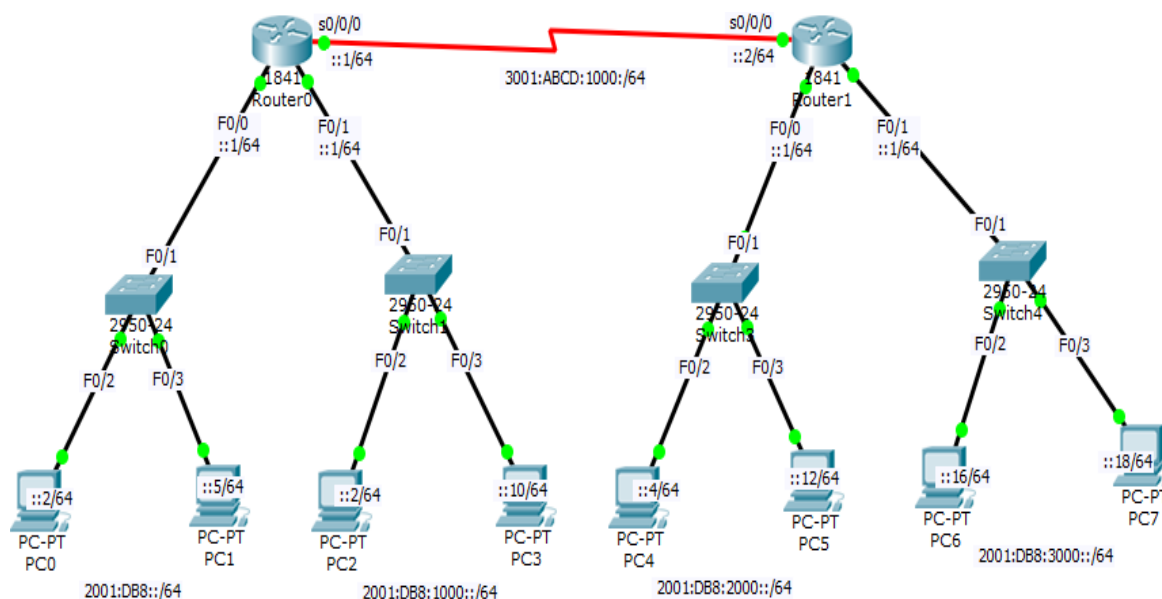
Hardware	Software
Computadora con los siguientes requisitos: ➤ Procesador mínimo de velocidad de 2.1 GHz ➤ Memoria RAM de 1 GB.	Simulador de redes Packet Tracer 6.1.1 con los siguientes elementos: ➤ 2 Router cisco de la serie 1841 ➤ 4 Switches 2950-24 ➤ 8 PCs

Conocimientos previos

Para la correcta realización de esta práctica el estudiante deberá tener conocimientos básicos de direccionamiento IPv6, Puertas de enlace (Gateway), dirección de red IPv6 y enrutamiento estático IPv6.



Topología



Funcionalidad

En la figura de esta práctica se muestra la topología en la cual se estará asignando direcciones IPv6 a los equipos finales, y a su vez se estará configurando rutas estáticas IPv6 en los routers para la comunicación de los equipos finales pertenecientes a diferentes redes LAN.

Comandos de ayuda

Comando	Descripción
configure terminal	Entra en el modo de configuración global desde el modo EXEC privilegiado
Enable	Entra en el modo EXEC privilegiado
End	Finaliza y sale del modo de configuración
ipv6 unicast-routing	Habilita el enrutamiento IPv6 que viene desactivado por defecto.
ipv6 address <IPv6-prefix/prefix-length>	Asigna una dirección IPv6 a una interfaz.
interfaz Tipo Número	Cambios en el modo de configuración global al modo de configuración de interfaz



ipv6 route ipv6-prefix/prefix-length {ipv6-address interface-type interface-number [ipv6-address]} [administrative-distance] administrative-multicast distance unicast multicast] [tag tag]	Establece una ruta estática en IPv6.
no shutdown	Activa la interfaz
ping IPv6-Address	Envía un ICMP echo Request to the specified address

Datos de los routers

Dispositivos	Interfaz	Dirección IP	Dirección de Red
Router0	F0/0	2001:DB8::1/64	2001:DB8::/64
	F0/1	2001:DB8:1000::1/64	2001:DB8:1000::/64
	S0/0/0	3001:ABCD:1000::1/64	3001:ABCD:1000::/64
Router1	F0/0	2001:DB8:2000::1/64	2001:DB8:2000::/64
	F0/1	2001:DB8:3000::1/64	2001:DB8:3000::/64
	S0/0/0	3001:ABCD:1000::2/64	3001:ABCD:1000::/64

ENUNCIADO

Asignación de direcciones IPv6

Establecer las direcciones IPv6 a las interfaces de los equipos que se mencionan a continuación y asignar la dirección de su puerta de enlace (gateway) si el equipo lo requiere, rellenando el siguiente cuadro a como se muestra en la topología anterior:

Dispositivos	Dirección IP	Máscara de red	Gateway
PC0	2001:DB8::2	/64	2001:DB8::1
PC1			
PC2	2001:DB8:1000::2	/64	2001:DB8:1000::1
PC3			
PC4	2001:DB8:2000::4	/64	2001:DB8:2000::1
PC5			
PC6	2001:DB8:3000::16	/64	2001:DB8:3000::1



PC7			
-----	--	--	--

Pruebas de comunicación entre equipos finales

Habiendo asignado correctamente las direcciones IPv6 y Gateway a las PCs y habiendo asignado la dirección a los routers haremos la prueba de comunicación entre estos dispositivos mediante el comando Ping.

Hacer ping desde PC0 a la dirección del Router0.

Hacer ping desde PC0 a la dirección del PC4.

Hacer ping desde PC5 a la dirección del Router1.

Hacer ping desde PC5 a la dirección del PC7

Enrutamiento estático IPv6 (Solución a problemas de conectividad)

Cómo pudimos observar al probar la comunicación de los equipos PC0 y PC4 usando el comando ping no lo hacían debido a que los dispositivos están en diferentes redes LAN separados por routers, por ende se buscó una alternativa para su correcta comunicación entre ellos. Así que se hace uso de enrutamiento estático para solventar ese problema.

Router0

```
Router0(config)#IPv6 route 2001:DB8:2000::/64 3001:ABCD:1000::2  
Router0(config)#IPv6 route 2001:DB8:3000::/64 3001:ABCD:1000::2
```

Router1:

```
Router1(config)#IPv6 route 2001:DB8::/64 3001:ABCD:1000::1  
Router1(config)#IPv6 route 2001:DB8:1000::/64 3001:ABCD:1000::1
```

Tiempo estimado de solución

- 4 horas



Preguntas de análisis

1. De acuerdo a los conocimientos adquiridos ¿Cuáles son las ventajas que presenta IPv6 hacia IPv4?
2. ¿Es correcto asignar la dirección 2001:BACA:1000::/64 a la interfaz de un router? Justifique.
3. ¿Cuál es el motivo de que IPv6 sea un direccionamiento sin clases?
4. ¿Cuántos bits forman una dirección IPv6?
5. ¿Es correcto asignar la dirección 3001:DB8::/64 como puerta de enlace (Gateway)?
6. Es posible asignar una dirección link-local a una dirección física. Justifique Sí o No ¿Por qué?



PRÁCTICA 2: AUTOCONFIGURACIÓN Y DHCPv6.

Objetivo general

- Mostrar el funcionamiento de la asignación de direcciones ip mediante autoconfiguración y DHCPv6.

Objetivos específicos

- Establecer el funcionamiento de la autoconfiguración y la forma correcta de su implementación en el software de simulación de redes GNS3.
- Asignar direcciones de manera dinámica mediante DHCPv6.

Introducción

En el siguiente laboratorio práctico se estudiara el funcionamiento del protocolo DHCP que permite distribuir direcciones ip de manera dinámica, al igual que el de autoconfiguración para la distribución de direcciones de forma automática.

Requerimientos:

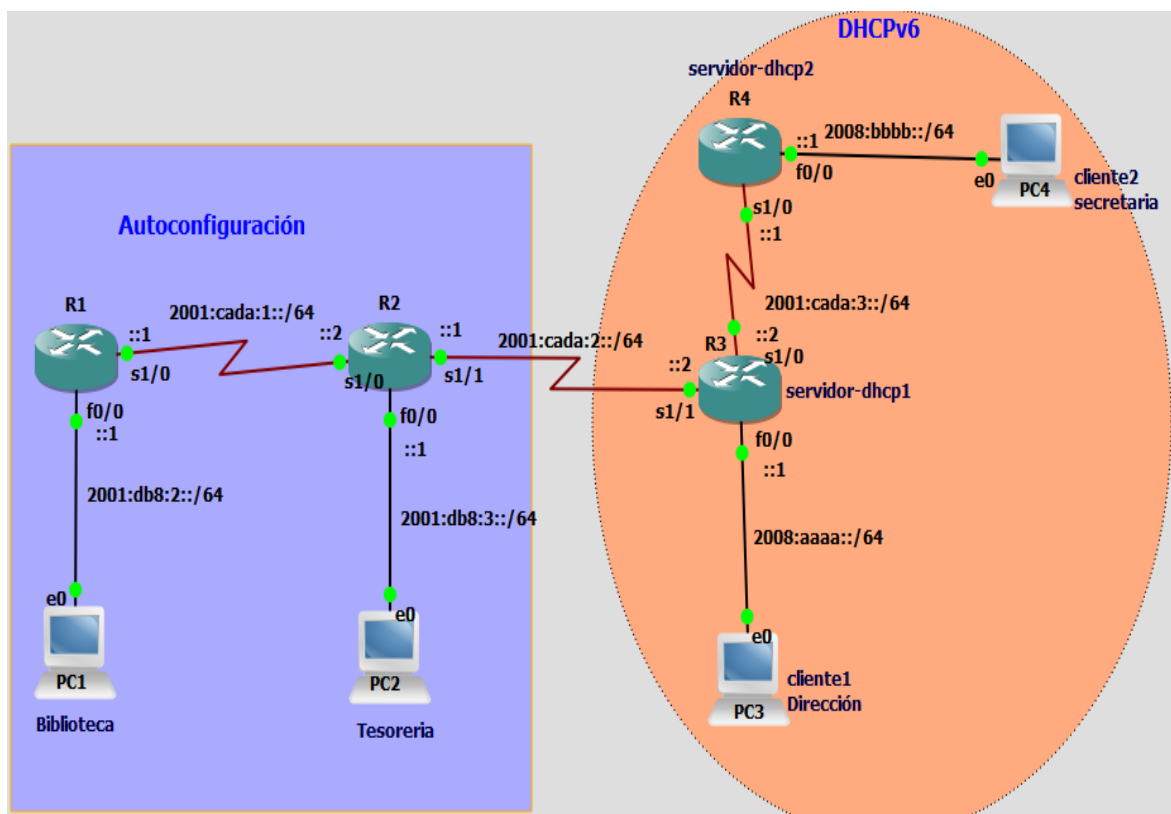
Hardware	Software
Computadora con los siguientes requisitos: ➤ Procesador mínimo de velocidad de 2.1 GHz ➤ Memoria RAM de 2 GB.	Simulador GNS3 con los siguientes elementos: ➤ 4 Routers serie C7200. ➤ 4 PCs.

Conocimientos previos

Para la correcta realización de esta práctica es necesario que el alumno tenga conocimientos previos de DHCPv6 y Autoconfiguración que se adoptaron en la parte teórica de este documento. Recuerde que la asignación de direcciones ip se puede realizar de 2 maneras: ya sea de forma automática (Autoconfiguración) o de forma dinámica (DHCPv6).



Topología



Funcionamiento

En la topología anterior se pretende implementar DHCPv6 en el cual se creará un pool de direcciones ip para asignar direcciones de manera dinámica a las PCs que están conectadas a los routers R1 y R3 respectivamente. El protocolo de Autoconfiguración se establecerá en los routers R5 y R6. Además de utilizar el enrutamiento estático IPv6 en los routers de este escenario para la comunicación entre las distintas redes IPv6.

Comandos de ayuda

Comando	Descripción
ipv6 dhcp pool name-pool	Crea el nuevo pool de direcciones.
ipv6 dhcp server name-pool	Define el pool del servidor dhcp.
ipv6 local pool name-pool address-ipv6 long-prefi	Define el rango de prefijo que se utilizara en la asignación dinámica.
ipv6 enable	Habilita la autoconfiguración.
ip auto	Solicita la dirección ip desde las pc.



sh ipv6 dhcp pool nom-pool	Visualiza la información del pool especificado.
sh ipv6 all	Muestra la dirección de una VPcs.

Datos de los dispositivos

AUTOCONFIGURACIÓN		
R1		
Interfaz	Dirección de red	Gateway
Fa0/0	2001:db8:2::/64	2001:db8:2::1/64
S1/0	2001:cada:1::/64	2001:cada:1::1/64
R2		
Fa0/0	2001:db8:3::/64	2001:db8:3::1/64
S1/0	2001:cada:1::/64	2001:cada:1::2/64
S1/1	2001:cada:2::/64	2001:cada:2::1/64
Datos de las PC		
Nombre	Interfaz del router	Dirección de Red
PC1: Biblioteca	Fa0/0	2001:db8:2::/64
PC2: Tesorería	Fa0/0	2001:db8:3::/64

DHCPv6			
R3: Servidor-dhcp1			
Interfaz	Pool-nombre	Dirección de red	Gateway
Fa0/0	Dirección	2008:aaaa::/64	2008:aaaa::1/64
Se1/0	-----	2001:cada:3::/64	2001:cada:3::2/64
R4: Servidor-dhcp2			
Fa0/0	Secretaria	2008:bbbb::/64	2008:bbbb::1/64
Se1/0	-----	2001:cada:3::/64	2001:cada:3::2/64
S1/1	-----	2001:cada:2::/64	2001:cada:2::2/64
Dato de las PC			



Nombre	Interfaz del router	Dirección de Red
PC3: cliente1 Dirección	Fa0/0	2008:aaaa::/64
PC4: cliente2 Secretaria	Fa0/0	2008:bbbb::/64

ENUNCIADO

En esta práctica se pretende configurar una topología de red como la mostrada anteriormente. Se deberá realizar lo siguiente:

- En R1 y R2, se llevara a cabo la implementación de Autoconfiguración, tomando como base las direcciones de red de la tabla **Autoconfiguración**.
- En R3 y R4, se asignara direcciones IP por DHCP, en los cuales se crearan Pools de direcciones en cada router a como se muestra en la tabla **DHCPv6**.
- Una vez que se configuraron los routers verificar las direcciones IP en cada una PCs.

NOTA: El comando ip auto se usa para la solicitud de direcciones IP dinámicas usando el protocolo DHCPv6. Definir el enrutamiento estático para la comunicación de toda la red.

Tiempo estimado de solución.

- 5 horas

Preguntas de análisis

1. ¿Cuál es la diferencia entre DHCPv6 y autoconfiguración?
2. ¿Un router puede recibir dirección ip por autoconfiguración? Justifique su respuesta.
3. ¿Qué comando se debe incluir para que la autoconfiguración se lleve a cabo?
 - a. #ipv6 unicast-routing
 - b. #ipv6 enable
 - c. #ipv6 address address-ipv6-prefix



4. ¿Un pool de direcciones ipv6 se crea con el siguiente comando?
 - a. #ipv6 dhcp pool name-pool
 - b. #ip dhcp pool name-pool
 - c. #ipv6 dhcp server name-pool

5. ¿Qué comando necesita las PC para solicitar su dirección por DHCP?
 - a. #show ipv6 all
 - b. #show ipv6
 - c. ip auto



PRÁCTICA 3: MECANISMOS DE TRANSICIÓN IPv4 E IPv6

Objetivo general

- Identificar y definir los mecanismos de transición que se pueden implementar entre IPv4 e IPv6.

Objetivos específicos

- Seleccionar que mecanismos de traducciones de direcciones IPv4 e IPv6 (Túneles de Transición, Dual Stack) se aplicaran en el laboratorio.
- Configurar los mecanismos de transición seleccionado para su correcta implementación.
- Verificar la conectividad entre equipos y señalar que los mecanismos de traducciones seleccionadas son viables.

Introducción

Un aspecto muy importante en el desarrollo del protocolo IPv6 es lograr la transición de IPv4 a IPv6, debido a que ambos no son compatibles. IPv6 ofrece mejores características que IPv4, no obstante debemos tener en cuenta que el cambio a este nuevo protocolo es un proceso muy alentador pero que a la vez es lento, debido a los altos costos de equipos y sobre todo a que el protocolo actual todavía está en funcionamiento.

Desde la creación del nuevo protocolo surgieron los siguientes debates: cómo lograr la comunicación de IPv6 dentro de IPv4 y cómo establecer el tipo de comunicación, es de ahí que nacen los mecanismos de transición, en la actualidad se implementan los siguientes mecanismos básicos:

- **Dual Stack:** proporciona soporte para manejar simultáneamente el protocolo IPv4 e IPv6. En una red dual Stack, ambos protocolos son desplegados completamente y los protocolos de enrutamiento deben llevar los prefijos correspondientes a cada tecnología, de manera transparente.
- **Túneles:** permite a los equipos que implementan el protocolo IPv6 comunicarse con otras redes IPv6 sobre la infraestructura de IPv4 y viceversa, existen diferentes tipos de túneles, entre los cuales tenemos:
 - **Túneles Manuales:** tal como su nombre lo indica, se configuran en forma manual tanto en un extremo como en el otro del túnel. Esta solución, si bien funciona, impone establecer el túnel de forma estática con algún dispositivo remoto que pueda proveernos conexión hacia redes IPv6.
 - **Túneles Automáticos:** Al contrario de los manuales, no es necesario configurar en forma estática en ambos extremos sino que se establecen automáticamente con una configuración mínima.
 - **4in6:** Encapsula tráfico IPv4 en IPv6.
 - **6in4:** Encapsula tráfico IPv6 en IPv4.



- **6to4:** Permite tráfico IPv6 sobre una red IPv4 sin la necesidad de configurar túneles de forma explícita, aunque se mantiene la función de encapsulamiento de IPv6 en IPv4.
- **NAT64:** Es un mecanismo que permite a hosts IPv6 comunicarse con servidores IPv4.

La implementación del nuevo protocolo de internet es un avance en la comunicación de dispositivos que trabaja con el servicio de internet.

La realización de este laboratorio le permitirá al estudiante comprender de manera correcta y eficiente como trabajar los siguientes mecanismos de transición: **Dual Stack y Túneles Manuales.**

Requerimientos:

Hardware	Software
➤ Computadora con los siguientes requisitos: • Procesador mínimo de velocidad de 2.1 GHz • Memoria RAM de 1 GB.	➤ Simulador Packet Tracer 6.1.1 : • 10 Routers • 9 Pc • Servidor DHCP (IPv4).

Conocimientos previos

Para la correcta realización de esta práctica debe tener conocimientos:

- Asignación de direcciones IPv6 e IPv4.
- Configuración de enrutamiento IPv6 e IPv4.
- Configuración servidores DHCP para IPv4 e IPv6.

Topología





Funcionalidad

En la imagen anterior se observa una topología de red, en la cual se aplicaron los dos mecanismos de transición antes mencionados.

La topología está estructurado de la siguiente manera:

Sistema Autónomo 1: Se aplicó el mecanismo de transición conocido como Dual Stack, por lo cual IPv6 e IPv4 se ejecutan de manera simultánea. En IPv6 se configuro el enrutamiento estático, mientras que en IPv4 se aplicó enrutamiento dinámico conocido como: RIPv2, los equipos finales se le asigno direcciones IP de manera estática tanto para IPv6 e IPv4.

Sistema Autónomo 2: Está estructurado en una red de direcciones IPv4, consta de un servidor DHCP, el cual asigna direcciones dinámicas a los equipos finales y el protocolo de enrutamiento dinámico establecido para este sistema fue RIPv2, lo cual permitirá establecer conexión con los equipos que implemente IPv4 en el sistema Autónomo 1.

Sistema Autónomo 3: Está estructurado en una red de direcciones IPv6, se aplicó enrutamiento estático y los equipos finales se le configuro tanto direcciones estática como dinámicas.

NOTA: Entre los sistemas Autónomos 2 y 3 se crearon túneles manuales, específicamente en los siguientes routers:

- Tunel0: Abarca los routers0 y routers4.
- Tunel1: Abarca los routers0 y routers1.

Comandos Generales:

Comando o Acción	Descripción
enable	Habilita el modo EXEC privilegiado.
configure terminal:	Entra en el modo de configuración global.
ipv6 unicast-routing	Permite el envío de datagramas IPv6 unicast.
interface type number	Configura un tipo de interfaz y entra en la interfaz el modo de configuración
ipv6 enable	Permite el procesamiento de IPv6 en una interfaz.
ipv6 address {ipv6-address/prefix-length prefix-name sub-bits/prefix-length}	Configura una dirección IPv6 en base a un prefijo y permite el procesamiento de IPv6 en una interfaz.
DHCPv6	



ipv6 dhcp pool pool-name Example: Router(config)# ipv6 dhcp pool pool1	Configura un conjunto de información de configuración DHCPv6 y entra en el modo de configuración de la agrupación DHCPv6.
prefix-delegation pool pool-name [lifetime valid-lifetime preferred-lifetime] Example: Router(config-dhcp)# prefix-delegation pool pool1 lifetime 1800 60	Especifica un prefijo IPv6 pool local llamado desde la que los prefijos se delegan a los clientes DHCPv6.
Exit	Exits DHCPv6 pool configuration mode configuration mode, and returns the router to global configuration mode.
interface type number	Especifica un tipo de interfaz y el número, y coloca el router en el modo de configuración de interfaz.
ipv6 dhcp server pool-name [rapid-commit] [preference value] [allow-hint] Example: Router(config-if)# ipv6 dhcp server pool1	Enables DHCPv6 on an interface.
Router(config)#ipv6 local pool pool1 2001:4000::/40 64 Example: ipv6 local pool pool1 2001:4000::/40 64	Define un pool de bloque de direcciones que se le asignaran al cliente.
Enrutamiento Estático	
ipv6 route ipv6-prefix/prefix-length {ipv6-address interface-type interface-number [ipv6-address]} [administrative-distance] administrative-multicast distance unicast multicast] [tag tag]	Establece una ruta estática en IPv6.
no shutdown	Activa la interfaz
Túnel Manual	
Router(config)# interface tunnel number	Crea una interfaz de túnel que es virtual.
Router(config-if)# tunnel source { type number}	Configura el origen del túnel.
Router(config-if)# tunnel destination { ip-address}	Comando de configuración de interfaz que especifica la dirección de destino para una interfaz de túnel. En este caso el parámetro ip-address es una dirección IPv4.



Router(config-if)# tunnel mode ipv6ip

Comando de configuración que permite la encapsulación de paquetes IPv6 en IPv4.

Direccionamiento Dual Stack / Sistema Autónomo 1

Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
Routers8	172.28.2.1/30	192.168.20.2/30	190.45.12.1/30		RIPv2
	2001:DB:AA::1/64	2010:ABC:DB8:1::5/64	2010:ABC:DB8:1::/64		Estático
Routers9	172.28.2.2/30	192.168.50.1/24			RIPv2
	2001:DB:AA::2/64			2001:D:A::1/64	Estático
Routers10	10.17.0.1/24	192.168.20.1/30			RIPv2
		2010:ABC:DB8:1::4/64		2004:AC:D B::1/64	Estático
Routers1			190.45.12.2/30		RIPv2
			2009:AA:11::20/64		Estático

Direccionamiento IPv4 / Sistema Autónomo 2

Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
Routers0		192.168.1.1/30			RIPv2
Routers1	192.1.20.1/30	192.168.1.2/30		120.10.40.1/30	RIPv2
Routers2	192.1.20.2/30	192.168.10.1/30			RIPv2
Routers3	192.16.5.1/30	192.168.10.2/30			RIPv2
Routers4	192.16.5.2/30				RIPv2
Routers7		150.240.20.0/24	150.240.20.1/24	120.10.40.2/30	RIPv2

Direccionamiento IPv6 / Sistema Autónomo 3



Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
Routers0	2002:3:3::1/64				Estático
Routers4		2002:2:2::1/64			Estático
Routers5	2019::1/64	2002:2:2::2/64			Estático
Routers6	2002:3:3::2/64	2015::1/64			Estático

Pools				
IPv6/IPv4				
Pool Name	Gateway	Dirección	Subredes	Inicio Dirección IP
Cliente	2019::1/64	2019::/40	24	2019::2/64
Administración	140.240.80.1/24	140.240.80.0/24		140.240.80.2/24

Servidores DHCP				
Equipo	Interfaz	Subinterfaces	Dirección IP	Gateway
Router5	F 0/0		2019::1/64	2019::1/64
DHCP	F 1/0		150.240.20.2/24	150.240.20.1/24

Equipos de Escritorios / Laptops /Servidor					
Nombre	Interfaz	Dinámico	Estático	Autoconfiguración	Dirección IPv6 e IPv4
Laptop1	Fa 0	Si			2019::1/64
Laptop2					
Laptop3	Fa 0		Si		10.17.0.10
Laptop4	Fa 0		Si		2004:AC:DB::10
Laptop5	Fa 0	Si			140.240.80.0/24
Laptop6					
Laptop7	Fa 0		Si		2001:D:A::2
Laptop8	Fa 0		Si		192.168.50.10
Laptop9	Fa 0		Si		2015::2
DHCP	Fa 0		Si		150.240.20.2



Túnel					
Routers	interface Túnel	Ipv6 Address	Túnel Source	Túnel Destination	túnel Mode
Router0	Tunnel0	2001:DB8:1:1::1/64	FastEthernet0/1	192.16.5.2	ipv6ip
Router4		2001:DB8:1:1::2/64	FastEthernet0/0	192.168.1.1	
Router1	Tunnel1	2016::1/64	FastEthernet0/1	192.168.1.1	ipv6ip
Router0		2016::2/64	FastEthernet0/1	192.168.1.2	

ENUNCIADO

Asignación de direcciones IP.

Se deberá asignar las direcciones IPv4/IPv6 a las interfaces de los routers, tal y como aparecen en el cuadro llamado: **Direccionamiento Dual Stack / Sistema Autónomo 1, Direccionamiento IPv4 / Sistema Autónomo 2 y Direccionamiento IPv6 / Sistema Autónomo 3.**

Asignación de direcciones estática a equipos.

Se deberá asignar su respectiva dirección estática a los equipos como laptops y servidores como se visualiza en el cuadro llamado: **Equipos de Escritorios / Laptops /Servidor.**

Configuración de enrutamiento Estático en redes IPv6.

Para lograr la conectividad entre las diferentes redes IPv6, es necesario que cada router anuncie las redes que tienen directamente conectada, y una forma para llegar hacia rutas desconocidas, es aplicando enrutamiento estático para IPv6. (Ver cuadro de **Comando General**).

Configuración DHCP

Para la asignación de direcciones dinámicas se deberá crear los Pool tal y como se muestra en los siguientes cuadros: **Pools y Servidores DHCP.**

Configuración de Mecanismo de Transición

En la práctica se implementó mecanismo transición como Túneles Manuales y el Dual Stack con el propósito de mostrar la coexistencia que pueden tener IPv4 con IPv6.

- Para la configuración del túnel ver cuadro **Comando General** y para ver cómo se plasmaron los túneles creado en esta práctica ver el **cuadro Túnel.**
- Recordar que el Dual Stack hace trabajar paralelamente ambos protocolos (IPv6 e IPv4).



Tiempo estimado de solución

➤ 7 horas

Preguntas de Análisis.

En los siguientes enunciados se pretende que los estudiantes sean capaces de analizar y responder las siguientes preguntas con respecto al tema de mecanismos de transición IPv4/IPv6, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. ¿Qué entiende por mecanismos de transición?
2. ¿Cuáles son las ventajas y desventajas que presentan los mecanismos de transición?
3. En un túnel manual IPv4 ¿En qué consiste la encapsulación de paquetes IPv6 dentro de redes IPv4?
4. ¿Explique el proceso que se produce cuando se comunican 2 redes IPv6 usando un túnel manual IPv4?
5. ¿Qué es el mecanismo DUAL STACK? ¿Por qué es importante usar el mecanismo de transición DUAL STACK? Justifique su respuesta.
6. Indique la afirmación correcta en los siguientes enunciados.
 - a) En el mecanismo de transición todas las redes tanto IPv4 como IPv6 se comunican sin ningún tipo de problema.
 - b) Los protocolos de enrutamiento dinámico de IPv4 no se pueden aplicar porque tiene problemas de compatibilidad con redes IPv6.
 - c) El mecanismo DUAL STACK es necesario en escenarios de red que solamente tienen redes IPv6 configuradas.
 - d) Los mecanismos DUAL STACK y túnel manual son necesarios en una topología de red que tiene tanto redes IPv4 como IPv6.
 - e) Ninguna de las anteriores.



PRÁCTICA 4: VLANs ESTÁTICAS Y DINÁMICAS.

Objetivos generales

- Crear VLANs estáticas analizando el funcionamiento y forma de aplicarlas.
- Configurar VLANs dinámicas en un entorno de red usando múltiples switches.

Objetivos específicos

- Evaluar los conceptos teóricos de VLANs estáticas y dinámicas, subnetting IPv6, enrutamiento estático, túneles IPv6-IPv4 y DHCPv6 utilizando la tecnología de cisco.
- Analizar el funcionamiento e interoperabilidad de las VLANs modo dinámico
- Comprender el funcionamiento del protocolo VTP de propagación de las VLANs dinámicas.
- Configurar subinterfaces y enrutamiento estático para la comunicación entre diferentes VLANs.
- Activar puertos troncales en los correspondientes switches.

Introducción

En el siguiente laboratorio práctico se estudiara el funcionamiento de las VLANs en forma estática y dinámica. El estudiante tendrá la posibilidad de simular una red física creando las VLANs respectivamente de forma estática o dinámica. Las VLANs son importantes en una topología de red aportando mayor seguridad a la red ya que brindan la facilidad de agrupar dispositivos en diferentes segmentos, esto con el fin de restringir el acceso a información a usuarios que no pertenecen a la misma VLANs.

Requerimientos:

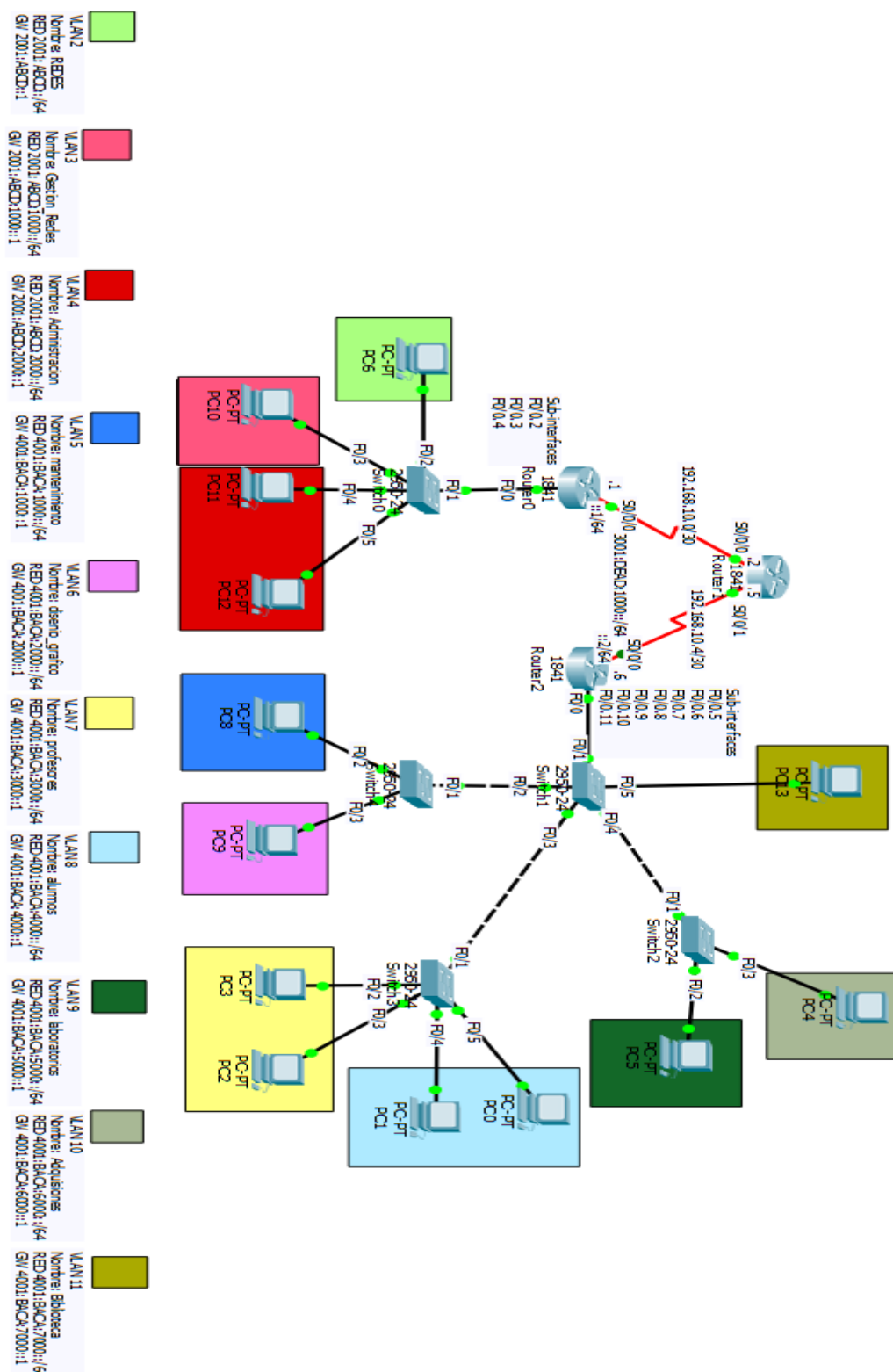
Hardware	Software
Computadora con los siguientes requisitos: ➤ Procesador mínimo de velocidad de 2.1 GHz ➤ Memoria RAM de 2 GB.	Simulador Cisco Packet Tracer 6.1.1 con los siguientes elementos: ➤ 3 Routers serie 1841. ➤ 5 Switch serie 2950-24. ➤ 13 Ordenadores.

Conocimientos previos

Para llevar a cabo la realización con éxito de la siguiente práctica, se necesita que el estudiante tenga conocimientos previos del establecimiento y funcionamiento tanto de VLANs estáticas como dinámicas e implementarlos en esta práctica para la correcta comprensión y funcionalidad de la misma.



Topología





Funcionalidad

En la topología anterior se pretende configurar tanto VLANs dinámicas como estáticas, a la vez que se estarán implementando de manera conjunta diferentes tecnologías y protocolos a como se mencionan a continuación:

- Se crearán 4 VLANs de forma estática y 6 de manera dinámicas con sus respectivas configuraciones las cuales estarán configuradas en los switches 0 y 1 respectivamente.
- Habrá al menos un protocolo de enrutamiento, en nuestro caso enrutamiento estático para la comunicación entre las diferentes subredes pertenecientes a las distintas redes LAN.
- Se hará uso de tecnologías túnel IPv6-IPv4 para la propagación de redes IPv6 dentro de redes IPv4.
- En las interfaces físicas Fa0/0 de los routers se estarán implementando sub-interfaces, esto con el fin de que el tráfico generado en un extremo de la red sea encaminado por los routers hacia el otro extremo, permitiendo así la comunicación entre las diferentes VLANs.
- Se deberá configurar los routers (Router0 y Router2) para brindar servicio DHCPv6 para la asignación de direcciones de manera dinámica a los equipos finales pertenecientes a las diferentes VLANs.

Comandos de ayuda

Comando	Descripción
vlan vlan-id	Crea una vlan
encapsulation dot1q vlan-id	Establece el método de encapsulación de la interfaz de enlace troncal 802.1Q VLAN, también especifica el ID de VLAN para la que las tramas deben ser etiquetados
ipv6 unicast-routing	Habilita el enrutamiento IPv6 que viene desactivado por defecto.
ipv6 address <IPv6-prefix/prefix-length>	Asigna una dirección IPv6 a una interfaz.
interface range fast ethernet Inicio puerto-finalización del puerto	Configura un rango de interfaces
interface Tipo Número	Cambios en el modo de configuración global al modo de configuración de interfaz
ip route destination-pre x destination-pre x mask {ip-address interface-type [ip-address]}	Establece una ruta estática
ipv6 route ipv6-prefix/prefix-length {ipv6-address interface-type interface-number [ipv6-	Establece una ruta estática en IPv6.



address]][administrative-distance] [administrative-multicast distance unicast multicast] [tag tag]	
no shutdown	Habilita una interfaz
show vlan	Muestra información de VLAN
switchport access vlan vlan-id	Establece la prioridad de árbol de expansión para su uso en el ID de puente
switchport access vlan vlan-id	Asigna la VLAN por defecto para un puerto
switchport mode {access dynamic {auto desirable} trunk}	Configura el modo de pertenencia a la VLAN de un puerto
vtp domain {nombre dominio}	Establece un dominio VTP
show vlan brief	Muestra las VLANs existentes
show vtp status	Muestra el estado de VTP
vlan database	Entra a la base de datos de las VLANs

Datos de los dispositivos

- Configuración VLANs estáticas

Switch 0 Switch 1			
Nombre Switch	Número de VLANs	Nombre	Puerto Access
Switch 0	2	REDES	F 0/2
	3	Gestión_Redes	F 0/3
	4	Administración	F 0/4, F 0/5
Switch 1	11	Biblioteca	F 0/5

- Configuración VLANs dinámicas

Switch 1			
Nombres de VLANs	Número de VLANs	Interfaz	Dirección de Red
Mantenimiento	5	Vlan 5	4001:BACA:1000::/64



Disenio_gráfico	6	Vlan 6	4001:BACA:2000::/64
Profesores	7	Vlan 7	4001:BACA:3000::/64
Alumnos	8	Vlan 8	4001:BACA:4000::/64
Laboratorios	9	Vlan 9	4001:BACA:5000::/64
Adquisiciones	10	Vlan 10	4001:BACA:6000::/64

Switches		
Nombre	Número de VLANs	Puertos Access
Switch 4	5	F 0/2
Switch 4	6	F 0/3
Switch 3	7	F 0/2, F0/3
Switch 3	8	F 0/4, F 0/5
Switch 2	9	F 0/2
Switch 2	10	F 0/3

- Routers

Routers				
Equipos	Interfaz	Subinterfaces	Dirección IPv6	Etiquetado de Vlan
Router 0	F 0/0	F 0/0.2	2001:ABCD::1/64	2
		F 0/0.3	2001:ABCD:1000::1/64	3
		F 0/0.4	2001:ABCD:2000::1/64	4
Router 2	F 0/0	F 0/0.5	4001:BACA:1000::1/64	5
		F 0/0.6	4001:BACA:2000::1/64	6
		F 0/0.7	4001:BACA:3000::1/64	7
		F 0/0.8	4001:BACA:4000::1/64	8
		F 0/0.9	4001:BACA:5000::1/64	9
		F 0/0.10	4001:BACA:6000::1/64	10
		F 0/0.11	4001:BACA:7000::1/64	11



POOLs				
Pool Name	Gateway	Dirección IPv6	Subredes	Inicio Dirección IPv6
REDES	2001:ABCD::1	2001:ABCD::/40	24	2001:ABCD::/64
Gestion_Redes	2001:ABCD:1000::1	2001:ABCD:1000::/40	24	2001:ABCD:1000::/64
Administracion	2001:ABCD:2000::1	2001:ABCD:2000::/40	24	2001:ABCD:2000::/64
Mantenimiento	4001:BACA:1000::1	4001:BACA:1000::/40	24	4001:BACA:1000::/64
disenio_grafico	4001:BACA:2000::1	4001:BACA:2000::/40	24	4001:BACA:2000::/64
Profesores	4001:BACA:3000::1	4001:BACA:3000::/40	24	4001:BACA:3000::/64
Alumnos	4001:BACA:4000::1	4001:BACA:4000::/40	24	4001:BACA:4000::/64
Laboratorios	4001:BACA:5000::1	4001:BACA:5000::/40	24	4001:BACA:5000::/64
Adquisiciones	4001:BACA:6000::1	4001:BACA:6000::/40	24	4001:BACA:6000::/64
Biblioteca	4001:BACA:7000::1	4001:BACA:7000::/40	24	4001:BACA:7000::/64

Servidores DHCPv6			
Equipos	Subinterfaces	Dirección IPv6	Gateway
Router 0	F 0/0.2	2001:ABCD::1/64	2001:ABCD::1/64
	F 0/0.3	2001:ABCD:1000::1/64	2001:ABCD:1000::1/64
	F 0/0.4	2001:ABCD:2000::1/64	2001:ABCD:2000::1/64
Router 2	F 0/0.5	4001:BACA:1000::1/64	4001:BACA:1000::1/64
	F 0/0.6	4001:BACA:2000::1/64	4001:BACA:2000::1/64
	F 0/0.7	4001:BACA:3000::1/64	4001:BACA:3000::1/64
	F 0/0.8	4001:BACA:4000::1/64	4001:BACA:4000::1/64



	F 0/0.9	4001:BACA:5000::1/64	4001:BACA:5000::1/64
	F 0/0.10	4001:BACA:6000::1/64	4001:BACA:6000::1/64
	F 0/0.11	4001:BACA:7000::1/64	4001:BACA:7000::1/64

ENUNCIADO

En esta práctica se pretende configurar una topología de red como la mostrada anteriormente. Se deberá seguir paso a paso la configuración de las diferentes tecnologías que se implementara en este escenario.

Creación de VLANs estáticas y dinámicas

➤ VLANs Estáticas

Como punto de partida, se deberán crear las VLANs en los dispositivos (switch 0, switch 1) cuyos datos se encuentran en el cuadro llamado **Switch0, Switch 1**.

Es en este punto en donde también se elegirán los puertos para los usuarios finales (puertos Access) perteneciente a la VLAN indicada en la configuración. En los cuadros anteriores también se muestra cuáles son los puertos que estarán operando en modo Access para el Switch 0, Switch 1.

➤ Propagación de VLANs dinámicas

En el equipo switch1 se crearán las VLANs con sus nombres y direcciones IPv6 a como se muestra en el cuadro llamado **switch1**.

En el Switch1 emplearemos el uso del protocolo VTP para la propagación de las VLANs, por lo que se deberá crear un dominio VTP llamado redes y su funcionalidad será modo server.

➤ Configuración de switches 2,3 y 4.

En los switch de capa 2 de la serie 2960-54 configurarlos a modo cliente VTP, de manera que estos también pertenezcan al dominio VTP que fue creado en el servidor VTP (switch1).

En los switches 2, 3 y 4 configurar los puertos en **modo access** a como se especifica en el cuadro **Switchs**, de modo que en estos equipos es donde se segmentan los puertos que permitirán el tráfico de las VLANs especificadas. Recuerde configurar estos equipos a modo cliente VTP, y que estos pertenezcan al dominio VTP creado en el switch 1.



➤ Servicio DHCPv6

En este escenario de red, se dará servicio DHCPv6 para las asignaciones de direcciones IPv6 de manera dinámica, esto con la finalidad de que todos los usuarios que pertenecen a cada una de las VLANs puedan obtener su dirección IPv6 sin necesidad de ingresarla de forma manual.

En los routers 0 y 2 se configurará los pools con el rango de direcciones que se les estarán asignando a los usuarios de cada VLANs como se especifica en la tabla **Pools**.

NOTA: El lifetime (tiempo de vida) del servicio DHCPv6 ofrecido por los routers 0 y 2 será igual a 3600. No olvide configurar las subinterfaces de los routers antes mencionados para que actúen como un server DHCPv6 con el comando **IPv6 DHCP server nombre (Nombre Pool)**. Las redes IPv6 para el uso de DHCPv6 de esta práctica en el router 0 y router 2 inicialmente tienen un prefijo /40 como se necesitan 24 subredes de esas redes el prefijo cambia a /64, **ejemplo: IPv6 local pool alumnos 4001:BACA:4000::/4064** y de allí se empieza a asignar direcciones de manera dinámica.

Túnel IPv6 e IPv4

El túnel para la comunicación de redes IPv6 debe de pasar por redes IPv4 que están configuradas en las interfaces seriales de los routers 0,1 y 2. El túnel IPv6/IPv4 se debe de configurar en los routers 0 y 2 donde la red que se usa es la 3001:DEAD:1000::/64.

Enrutamiento Estático

Para la comunicación de las distintas VLANs mostradas en este escenario se hará uso de enrutamiento estático tanto para redes IPv6 como para redes IPv4. El encaminamiento estático IPv6 se realizará en los routers 0 y 2 respectivamente al igual que el de enrutamiento estático IPv4.

NOTA: Cómo se puede observar en el escenario de red no hay configuradas redes IPv6 directamente sobre los enlaces seriales de los routers 0,1 y 2. Por tanto para que se puedan comunicar las VLANs usando Enrutamiento estático IPv6 se debe de usar las direcciones de red IPv6 del túnel anteriormente configurado.

Tiempo estimado de solución

- 5 horas



Preguntas de análisis

En los siguientes apartados se pretende que los estudiantes sean capaces de analizar y responder las siguientes preguntas con respecto al tema de VLANs, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. ¿Por qué la Vlan 1 no fue creada?
2. ¿Cuál es el objetivo de configurar el servidor VTP en el switch 1? ¿Por qué no se realiza de forma igual en los demás?
3. ¿Explique porque es importante la encapsulación do1qt en la creación de subinterfaces?
4. ¿Con respecto al tipo de VLANs configuradas en esta práctica, cual tipo de VLANs que puede ser el más factible implementar en una topología de red? .Justifique su respuesta.
5. ¿Cuáles son las ventajas de utilizar VTP en la topología de red anterior?
6. Encontró alguna diferencia para crear VLANs dinámicas o estáticas en IPv6 con respecto a IPv4.
¿Sí? ¿No? Explique.



PRÁCTICA 5: FRAME RELAY E INTERVLANS.

Objetivo general

- Implementar la tecnología Frame Relay y el enrutamiento entre VLANs usando el direccionamiento de red IPv6.

Objetivos específicos

- Crear subinterfaces en los routers que formen parte del escenario de red para que se comuniquen mediante Frame Relay.
- Implementar el protocolo NAT64 para la comunicación entre las VLANs IPv6 e IPv4.

Introducción

Con la implementación y realización de esta práctica se pretende que los alumnos sean capaces de entender y poder aplicar Frame Relay según su finalidad de uso. Además se hace uso del mecanismo de transición NAT64 para comunicar las diferentes redes IPv6 e IPv4 correspondientes

Requerimientos:

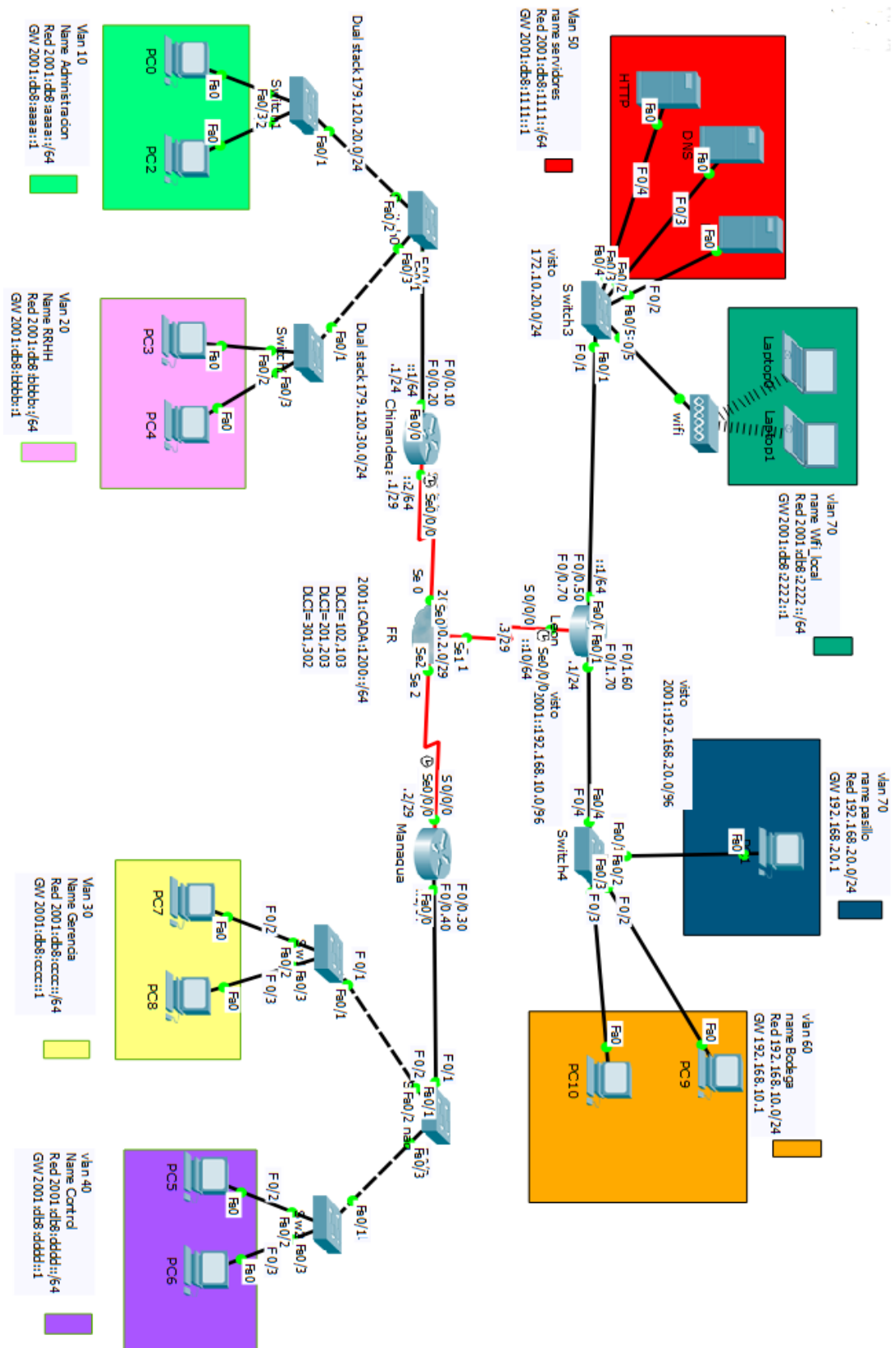
Hardware	Software
Computadora con los siguientes requisitos: ➤ Procesador mínimo de velocidad de 2.1 GHz ➤ Memoria RAM de 2 GB.	Simulador de redes Packet Tracer 6.1.1 con los siguientes elementos: ➤ 3 Router cisco de la serie 1841. ➤ 1 Nube PT. ➤ 8 Switches 2950-24. ➤ 13 PCs. ➤ 1 AccessPoint-PT ➤ 3 servidores.

Conocimientos previos

Para la correcta realización de esta práctica el alumno deberá tener conocimientos de implementación de Frame Relay, Intervlans y NAT64 ya que será de mucha importancia para realizar la misma. También conocimientos del uso y configuración de los servicios tales como DHCP (IPv6/IPv4), HTTP, DNS y FTP.



Topología.





Funcionalidad

Para el funcionamiento de este escenario de red se deberá de tomar en cuenta la implementación y configuración correcta de Frame Relay, InterVLANs y NAT64. Además de que se estarán usando de manera conjunta con diferentes tecnologías y protocolos a como se menciona a continuación:

- Se deberán crear ocho VLANs, asignándoles direcciones IP estáticas y dinámicas a los equipos finales perteneciente a cada una de ellas.
- Se crearán en la interfaz FastEthernet de los routers, sub-interfaces para comunicación entre VLANs distintas.
- Se asignara en la interfaz serial de los routers, los DLCI correspondientes para la comunicación en la nube frame relay.
- Habrá que configurar la nube Frame Relay con los DLCI que se configuró en las sub-interfaces de los routers.

Comandos de ayuda

Comando	Descripción
show frame-relay lmi {type number}	Muestra información sobre la interfaz de gestión local (LMI)
show frame-relay map.	Muestra las entradas del mapa frame relay actual e información acerca de las conexiones.
show frame-relay traffic	Muestra las estadísticas globales de frame relay desde la última recarga.
encapsulation frame-relay	Permite la encapsulación frame relay en la interfaz
frame relay map protocol protocol-address dlci	Define la correspondencia entre una dirección de protocolo de destino y el DLCI utilizado para conectarse a la dirección de destino.
frame-relay lmi-type {ansi cisco q933a}	Selecciona el tipo de interfaz de gestión local (LMI).

Datos de los dispositivos

- Configuración VLANs estáticas

Switches Leon1 y Leon2			
Nombre Switch	Número de VLANs	Nombre	Puerto Access
Leon1	50	servidores	F 0/2-4
	70	Wifi	F 0/5
Leon2	60	Bodega	F 0/2-3



	70	pasillo	Fa0/1
--	----	---------	-------

- Configuración VLANs dinámicas

Switches Chinandega1 y Managua1			
Nombres de VLANs	Número de VLANs	Interfaz	Dirección de Red
Administración	10	Vlan 10	2001:db8:aaaa::/64
RRHH	20	Vlan 20	2001:db8:bbbb::/64
Gerencia	30	Vlan 30	2001:db8:cccc::/64
Control	40	Vlan 40	2001:db8:dddd::/64

Switchs		
Nombre	Número de VLANs	Puertos Access
Switch 1	10	F 0/2-3
Switch 2	20	F 0/2-3
Switch 7	30	F 0/2-3
Switch 8	40	F 0/2-3

Routers					
Equipo	Subinterfaz	Dirección IP	Dirección de subred	Etiquetado VLANs	DLCI
Chinandega	Fa0/0.10	2001:db8:aaaa::1/64 179.120.20.1/24	2001:db8:aaaa::/64 179.120.20.0/24	10	
	Fa0/0.20	2001:db8:bbbb::1/64 179.120.30.1/24	2001:db8:bbbb::/64 179.120.30.0/24	20	
	Se0/0/0	2001:cada:1200::2/64	2001:cada:1200::/64		102,103
León	Fa0/0.50	2001:DB8:1111::1/64	2001:DB8:1111::/64	50	
	Fa0/0.70	2001:DB8:2222::1/64 192.168.40.1/24	2001:DB8:1111::/64 192.168.40.0/24	70	
	Fa0/1.60	192.168.10.1/24	192.168.10.0/24	60	
	Fa0/1.70	192.168.20.1/24	192.168.20.0/24	70	



	Serial0/0/0	2001:cada:1200::10/64	2001:cada:1200::/64		201,203
Managua	Fa0/0.30	2001:db8:cccc::1/64	2001:db8:cccc::/64	30	
	Fa0/0.40	2001:db8:dddd::1/64	2001:db8:dddd::/64	40	
	Se0/0/0	2001:cada:1200::12/64	2001:cada:1200::/64		301,302

Pools				
Pool Name	Gateway	Dirección IP	Subredes	Inicio Dirección IP
wifi_local	2001:db8:2222::1	2001:db8:2222::/48	16	2001:DB8:2222::/64
Wifi	192.168.40.1	192.168.40.0/24		192.168.40.2/24
Administración	2001:db8:aaaa::1	2001:db8:aaaa::/48	16	2001:db8:aaaa::/64
Admin	179.120.20.1	179.120.20.0/24		179.120.20.2/24
RRHH	2001:db8:bbbb::1	2001:db8:bbbb::/48	16	2001:db8:bbbb::/64
RRHHs	179.120.30.1	179.120.30.0/16	8	179.120.30.2/24

Servidores DHCP			
Equipos	Subinterfaces	Dirección IP	Gateway
León	F 0/0.70	2001:db8:2222::1/64	2001:db8:2222::1/64
	F 0/0.70	192.168.40.1/24	192.168.40.1/24
Chinandega	F 0/0.10	2001:db8:aaaa::1/64	2001:db8:aaaa::1/64
	F 0/0.10	179.120.20.1/24	179.120.20.1/24
	F 0/0.20	2001:db8:bbbb::1/64	2001:db8:bbbb::1/64
	F 0/0.20	179.120.30.1/24	179.120.30.1/24

Tabla de direcciones con su respectiva traducción

NAT-64			
Equipo	Subinterfaz	Dirección-original	Traducción
León	Fa0/0.50	2001:DB8:1111::1/64	172.10.20.1/24
	Fa0/1.60	192.168.10.1/24	2001::192.168.10.1/96
	Fa0/1.70	192.168.20.1/24	2001::192.168.20.1/96
	Fa0/0.30	2001:db8:cccc::1/64	180.12.16.1/24



Managua	Fa0/0.40	2001:db8:dddd::1/64	180.160.17.1/24
Server FTP		2001:db8:1111::4/64	172.10.20.4/24
Server DNS		2001:db8:1111::5/64	172.10.20.5/24
Server HTTP		2001:db8:1111::6/64	172.10.20.6/24
PC1		192.168.20.3/24	2001::192.168.20.3/96
PC5		2001:db8:dddd::3/64	180.160.17.3/24
PC6		2001:db8:dddd::4/64	180.160.17.4/24
PC7		2001:db8:cccc::3/64	180.12.16.3/24
PC8		2001:db8:cccc::4/64	180.12.16.4/24
PC9		192.168.10.3/24	2001::192.168.10.3/96
PC10		192.168.10.4/24	2001::192.168.10.4/96

ENUNCIADO

Asignación de direcciones IP

Asignar direcciones IP a las subinterfaces FastEthernet de los siguientes equipos ver cuadro de **Routers** (Chinandega, León, Managua). Verifique que las direcciones sean ingresadas correctamente. De igual manera establecer las direcciones IP para cada interfaz serial, al igual que su DLCI correspondiente a cada una. (Ver tabla **Routers**). Además de ingresar las direcciones de forma manual a los equipos que están en la tabla **NAT64** sus direcciones IP originales como son los servidores **DNS, HTTP, FTP y PCs**.

Creación de VLANs estáticas y dinámicas.

Se deberán crear las VLANs en los dispositivos switches (**leon1 y leon2**) cuyos datos se encuentran en el cuadro llamado **Switches leon1 y leon2**.

Es en este punto en donde también se elegirán los puertos para los usuarios finales (puertos access) perteneciente a la VLAN indicada en la configuración. En los cuadros anteriores también se muestra cuáles son los puertos que estarán operando en modo access para los switches **leon1 y leon2**

En los dispositivos **Chinandega1 y Managua1** se crearán las VLANs con sus nombres y direcciones IPv6 a como se muestra en el cuadro llamado **Chinandega1 y Managua1**.

En los Switches **Chinandega1 y Managua1** se usa el protocolo VTP para la propagación de las VLANs por lo que se deberá crear un dominio VTP llamado Chinandega (**Chinandega1**), Managua (**Managua1**) y que su funcionalidad será modo servers.



NOTA: Recuerde que la funcionalidad VTP de los dispositivos que están en el cuadro **Switchs** será modo cliente y asegúrese de dar los puertos de acceso adecuados para cada VLANs descrita en el cuadro mencionado anteriormente.

Servicios DHCP, DNS, HTTP y FTP.

En este escenario de red, se dará servicio DHCP para las asignaciones de direcciones IP de manera dinámica, esto con la finalidad de que todos los usuarios que pertenecen a cada una de las VLANs puedan obtener su dirección IP sin necesidad de ingresarla de forma manual.

En los routers León y Chinandega se configurará los pools con el rango de direcciones que se les estarán asignando a los usuarios de cada VLANs como se especifica en la tabla **Pools**.

Las direcciones IP para brindar el servicio de DNS serán 2001:DB8:1111::5 y 172.10.20.5 respectivamente.

Recordemos de configurar el accesspoint-pt llamado **wifi** con contraseña WPA2-PSK.

El nombre del sitio web se llamara **telemática.com** y el nombre de dominio para el servidor ftp se le va a asignar como **ftp.telematica.com**.

Traducción Nat64 para la coexistencia de IPv6 e IPv4.

Los Routers León y Managua estarán configurados de modo que los paquetes de la red IPv6 y la red IPv4 se puedan comunicar usando el protocolo NAT64. tal como se muestra en el cuadro llamado **NAT64**.

Protocolo de encaminamiento de tráfico

Se deberá crear rutas por defecto en los routers nombrado Chinandega, León y Managua.

1. Router Chinandega

Una ruta por defecto que permita el paso del tráfico con destino a la red de las VLANs 30,40(una por cada VLANs).

Una ruta por defecto que permita el paso del tráfico con destino a la red de las VLANs 50, 60, 70(una por cada VLANs).

2. Router León

Una ruta por defecto que permita el paso del tráfico con destino a la red de las VLANs 10, 20 (una por cada VLANs).



Una ruta por defecto que permita el paso del tráfico con destino a la red de las VLANs 30,40 (una por cada VLANs).

3. Router Managua

Una ruta por defecto que permita el paso del tráfico con destino a la red de las VLANs 10,20 (una por cada VLANs).

Una ruta por defecto que permita el paso del tráfico con destino a la red de las VLANs 50,60, 70 (una por cada VLANs).

Configuración de la nube Frame Relay.

1. Serial 0

En la pestaña config de la nube Frame Relay Configurar el puerto serial 0 conectado directamente con el router nombrado Chinandega agregando los DLCI configurado en este router (ver cuadro de **Router** en el equipo Chinandega). Recuerde configurar los LMI, en este caso serán de tipo Cisco.

2. Serial 1

En la pestaña config de la nube Frame Relay Configurar el puerto serial 1 conectado directamente con el router nombrado León agregando los DLCI configurado en este router (ver cuadro de **Router** en el equipo León), Recuerde que el LMI será de tipo Cisco.

3. Serial 2

En la pestaña config de la nube Frame Relay Configurar el puerto serial 2 conectado directamente con el router nombrado Managua agregando los DLCI configurado en este router (ver cuadro de **Router** en el equipo Managua). Recuerde que el LMI será de tipo Cisco.

Tiempo estimado de solución

➤ 6 horas



Preguntas de análisis

1. ¿Qué comando del router muestra las tablas de Frame Relay?
 - a. show frame-relay map
 - b. show frame-relay pvc.
 - c. show frame-relay lmi
2. ¿Qué comandos pertenecen a la configuración de frame relay?
 - a. frame-relay map ipv6 2001:CADA:1200::12 102
 - b. encapsulation frame-relay
 - c. int fast0/0.10
3. ¿Cuál de las afirmaciones es correcta?
 - a. Las vlan se crean desde los switch configurados con VTP modo cliente.
 - b. Las vlan se crean desde el router al que están conectado.
 - c. Las vlan se crean desde los switch configurados con VTP modo servidor.
4. ¿Cuál de los comandos se debe configurar para crear el enlace del switch con el router?
 - a. switchport mode access
 - b. switchport mode trunk
5. Una vez creadas las subinterfaces para cada vlan que comando debe configurarse a continuación.
 - a. ip address 192.168.10.1 255.255.255.0
 - b. encapsulation dot1Q 60



PRÁCTICA 6: CONTROL DE LISTAS DE ACCESO.

Objetivo general

- Establecer listas de accesos fusionándolas con distintos protocolos y tecnologías en IPv6.

Objetivos específicos

- Configurar listas de acceso para brindar seguridad a una red LAN en el protocolo de direccionamiento IPv6.
- Mostrar la importancia de las listas de acceso para permitir y denegar el tráfico en redes IPv6.

Introducción

Con la implementación y realización de esta práctica, se pretende que los estudiantes sean capaces de entender y poder aplicar listas de accesos (ACL IPv6) según su tipo y finalidad de uso, además de comprender el funcionamiento de las mismas para dar seguridad a los equipos.

Requerimientos

Hardware	Software:
Computadora con los siguientes requisitos: ➤ Procesador mínimo de velocidad de 2.1 GHz ➤ Memoria Ram de 1 GB	Simulador Cisco Packet Tracer 6.1.1 con los elementos siguientes: ➤ 5 Routers 2811 ➤ 3 Servidores ➤ 6 Switchs serie 2960-24TT ➤ 18 PCs ➤ 2 Access Point-PT

Conocimientos previos

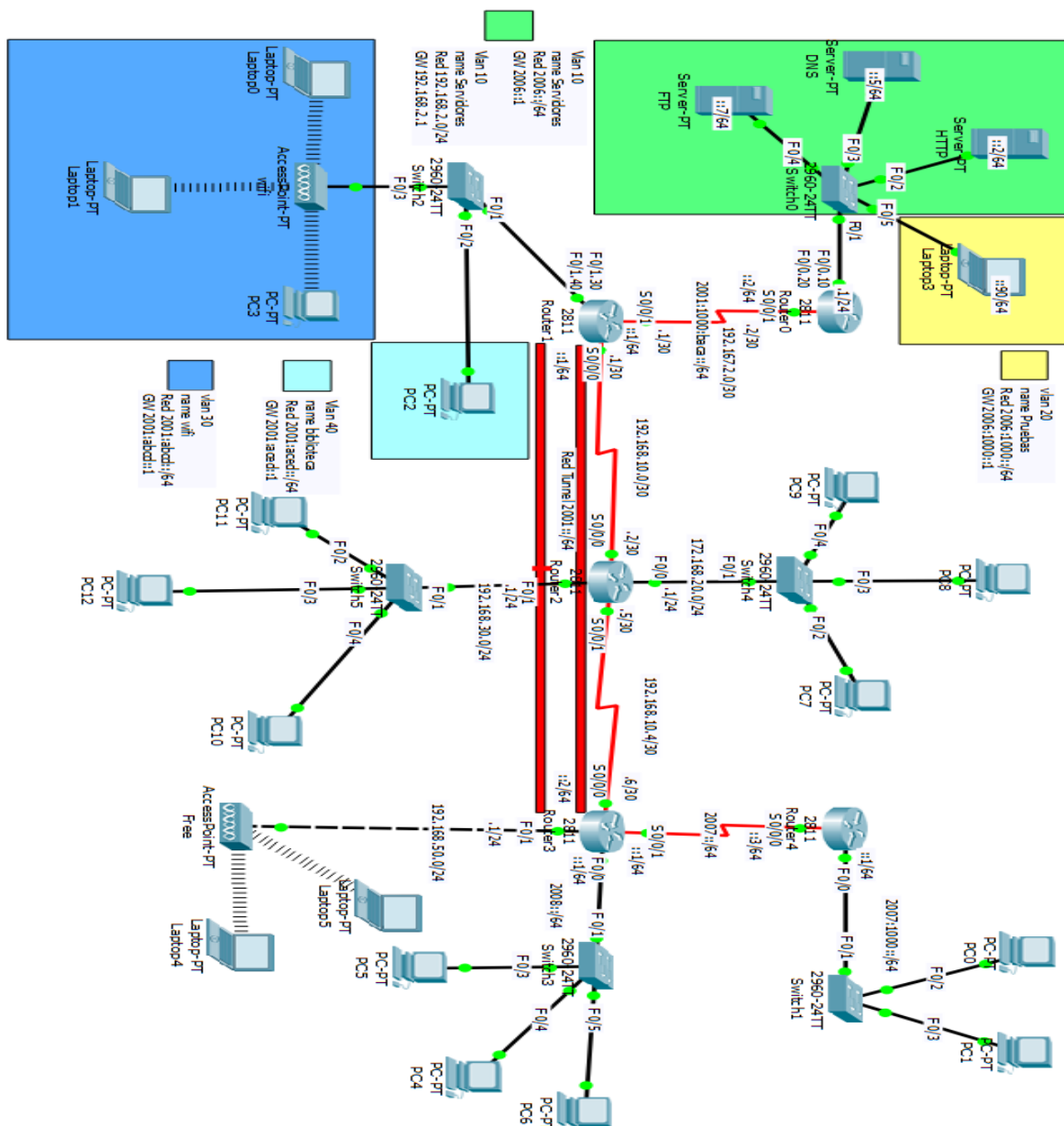
Para la correcta realización de esta práctica es necesario que el alumno tenga conocimientos básicos acerca del uso y funcionamiento de las listas de control de acceso (ACLs IPv6) proporcionado previamente en los aspectos teóricos, además conocimientos básicos de los siguientes temas:

- Direccionamiento IPv6/IPv4 (subnetting)
- Creación de Vlans estáticas en IPv6
- Asignación dinámica de direcciones IP mediante DHCP tanto en IPv4 como IPv6.
- Implementación de túnel para la comunicación entre redes IPv6 sobre redes IPv4.
- Enrutamiento estático en redes IPv6.
- Tipos de seguridad en redes inalámbricas WIFI.



NOTA: El protocolo de enrutamiento dinámico que se utilizó para la comunicación entre redes IPv4 es RIP versión 2.

Topología



Funcionalidad

En la figura de esta práctica se muestra la topología en la cual se estará implementando listas de accesos (ACLs IPv6) según la finalidad de uso, y a su vez se estará aplicando de manera conjunta con diferentes protocolos y tecnologías.



- Se deberán crear cuatro VLANs, donde se asignará direcciones de forma dinámica a los equipos finales perteneciente a cada una de ellas mediante el protocolo DHCPv6, agregando la dirección del Servidor DNS interno que poseerán los equipos de la red LAN.
- Se debe de configurar enrutamiento estático para la comunicación entre las distintas redes IPv6.
- Los equipos finales deberán acceder a cada uno de los servidores pertenecientes a la red LAN mediante DNS.
- El punto de acceso inalámbrico llamado **wifi** su protección será WPA2-PSK (contraseña **cisco12@**). Mientras que el que se llama Free será libre, es decir no tendrá ningún protocolo de seguridad inalámbrica.
- Se hace uso de túneles para la comunicación entre redes IPv6 sobre redes IPv4. Además el túnel se configurará en los routers Router1 y Router3.
- Por último como objetivo principal de la práctica, se deberán crear Listas de Control de Acceso (ACLs IPv6) para dar seguridad a cada uno de los equipos, manteniendo integridad y seguridad en el control de flujo de tráfico en la red LAN.

Comandos de ayuda

Comando	Descripción
<code>ipv6 unicast-routing</code>	Habilita el enrutamiento IPv6 que viene desactivado por defecto.
<code>ipv6 access-list access-list-name</code>	Define el nombre de la lista de acceso en ipv6. Además ingresa al modo de configuración
<code>deny permit protocol {source-ipv6-prefix/prefix-length any host source-ipv6-address} [operator [port-number]] {destination-ipv6-prefix/prefix-length any host destination-ipv6-address} [operator [port-number]] [dscp value] [fragments] [log] [log-input] [sequencevalue] [time-range name]</code>	Permite o niega un servicio en las listas de accesos en IPv6 (ACLs IPv6). Ejemplo: deny icmp 2001:ABCD::/64 host 2006::5 echo-request
<code>interface type number</code>	Establece la interfaz especificada y accede al modo de configuración
<code>ipv6 traffic-filter access-list-name {in out}</code>	Aplica la ACLs específica en la entrada o salida de tráfico en una interfaz.
<code>ipv6 access-class access-list-name { in out }</code>	establece quienes tienen acceso Vty usando ACLs



ip address [ip-address subnet-mask]	Asigna dirección IP a una interfaz
ipv6 address <IPv6-prefix/prefix-length>	Asigna dirección IPv6 a una interfaz
ip dhcp pool [nombre del pool]	Crea un POOL para asignar direcciones de manera Dinámica.
ipv6 dhcp pool [nombre del pool]	Crea un POOL para asignar direcciones de manera Dinámica en IPv6.
no shutdown	Habilita una Interfaz.
vlan database	Accede a la base de datos de VLANs
vlan [número-vlan name nombre-vlan]	crea una VLAN especificando el nombre y el número
encapsulation dot1q vlan-id	Establece el método de encapsulación de la interfaz de enlace troncal 802.1Q VLAN, también especifica el ID de VLAN para la que las tramas deben ser etiquetados
switchport access vlan vlan-id	Asigna la VLAN por defecto para un puerto
show vlan	muestra información de las VLANs

Datos de los dispositivos

VLANs			
Nombre Switch	Número de Vlan	Nombre	Puerto de acceso
Switch 0	10	Servidores	F 0/2-4
	20	Pruebas	F0/5
Switch 2	30	Wifi	F 0/3
	40	Biblioteca	F 0/2

Routers (subinterfaces)				
Equipos	Interfaz	Subinterfaces	Dirección IP	Etiquetado de Vlan
Router0	F 0/0	F 0/0.10	2006::1	10
			192.168.2.1	



		F 0/0.20	2006:1000::1	20
Router1	F 0/0	F 0/0.30	2001:abcd::1	30
		F 0/0.40	2001:aced::1	40

Routers				
Nombres	Direcciones IP			
	F 0/0	F 0/1	S 0/0/0	S 0/0/1
Router0	---	---	---	2001:1000:baca::2/64
Router1	---	---	192.168.10.1/30	2001:1000:baca::1/64
Router2	172.168.20.1/24	192.168.30.1/24	192.168.10.2/30	192.168.10.5/30
Router3	2008::1/64	192.168.50.1/24	192.168.10.6/30	2007::1/64
Router4	2007:1000::1/64	---	2007::3/64	---

POOLs				
Pool Name	Gateway	Dirección IP	Subredes	Inicio Dirección IP
Wifi	2001:abcd::1	2001:abcd::/48	16	2001:abcd::/64
Biblioteca	2001:aced::1	2001:aced::/48	16	2001:aced::/64
redipv4	172.168.20.1	172.168.20.0/16	8	172.168.20.2/24
redipv4.1	192.168.30.1	192.168.30.0/24	---	192.168.30.2/24
redipv4.2	192.168.50.1	192.168.50.0/24	---	192.168.50.2/64
redipv6	2008::1	2008::/48	16	2008::/64
redipv6.3	2007:1000::1	2007:1000::/48	16	2007:1000::/64

Servidores			
Nombres	Dirección IP	Gateway	Nombre de dominio
DNS	2006::5/64	2006::1	---
	192.168.2.5	192.168.2.1	



HTTP	2006::2/64 192.168.2.2	2006::1 192.168.2.1	ipv6acls.com <u>www.ipv6acls.com</u>
FTP	2006::7/64 192.168.2.7	2006::1 192.168.2.1	<u>ftp.ipv6acls.com</u>

ENUNCIADO

Asignación de direcciones IP

Como punto inicial de ésta práctica, empezaremos por asignarles direcciones IP a las interfaces de los dispositivos que se describen a continuación:

- Router0
- Router1
- Router2
- Router3
- Router4

Asignar direcciones IPv6 de manera estática a cada uno de los servidores tal y como se aprecia en la tabla llamada **servidores**.

NOTA: Recuerde establecer correctamente las direcciones IP para un funcionamiento satisfactorio. Los equipos que están conectados a la Vlan pruebas se le deben de ingresar su dirección IP incluyendo el servicio de DNS de forma manual.

Creación de VLANs

Crear las VLANs en los dispositivos Switch0 y 2 (ver cuadro **VLANs**), aquí también se elegirán los puertos para los equipos finales (puertos Access) perteneciente a la VLAN indicada en la configuración. En los cuadros anteriores también se muestra cuáles son los puertos que estarán operando en modo Access para el **Switch 0** y **Switch 2**.

NOTA: Recuerde que debe de asignar el **modo trunk** en las interfaces de los switches que crea conveniente.

Asignación de direcciones IP mediante DHCP a los equipos finales

En este escenario de red, se dará servicio DHCP tanto para IPv6 como IPv4 para las asignaciones de direcciones IP de manera dinámica y a su vez también se distribuirá el servicio de DNS, esto con la finalidad de que todos los usuarios que pertenecen a cada una de las redes puedan obtener su dirección IP sin necesidad de ingresarla de forma manual. Además de poder usar los servicios de red tales como HTTP y FTP.



En la mayoría de los Routers exceptuando Router0 se configurará los pools con el rango de direcciones que se les estarán asignando a los usuarios de cada una de las redes como se especifica en la tabla **Pools**. El servidor HTTP responderá a los nombres de dominios www.ipv6acls.com, ipv6acls.com tanto redes IPv6 como para IPv4 y el servidor FTP a su vez lo será para ftp.ipv6acls.com tal como aparece en el cuadro **Servidores**.

En todos los equipos finales (PC) se deberán configurar de forma que realicen una petición DHCP para que el pool que hemos creado asigne dirección de manera dinámica a cada uno de los equipos.

NOTA: El lifetime (tiempo de vida) del servicio DHCPv6 ofrecido por los routers1, 3 y 4 será igual a 3600. No olvide configurar las subinterfaces de los routers antes mencionados para que actúen como un server DHCPv6 con el comando **IPv6 DHCP server nombre (Nombre Pool)**. No nos olvidemos de los **APs** (Access point) con su respectiva configuración.

Configuración de protocolos de enrutamientos estático IPv6 y RIPv2 (IPv4).

Para encaminar el tráfico entre cada una de las redes existentes, se estará usando RIPv2 como protocolo de enrutamiento en las redes IPv4, Mientras que en las redes IPv6 se hará uso de enrutamiento estático. Los routers donde se hacen la configuración respectiva serán los siguientes:

Router1, Router2 y Router3: Aquí se configura **RIPv2** como protocolo de enrutamiento para la comunicación de redes IPv4.

Router0, Router1, Router3 y Router4: Se deberá de configurar enrutamiento estático IPv6 para la comunicación entre las distintas redes IPv6.

NOTA: En el túnel configurado en los Router1 y 3 también se hará uso de enrutamiento estático IPv6.

Implementación de las Listas de Control de Acceso (ACL IPv6)

Implemente las Listas de Control de Acceso creadas en el apartado anterior en los interfaces o subinterfaces de los dispositivos que crea conveniente y verifique su correcto funcionamiento realizando las debidas pruebas.

Tiempo estimado de solución

- 5 horas



Preguntas de Análisis

En los siguientes apartados se pretende que los estudiantes sean capaces de analizar y responder las siguientes preguntas con respecto al tema de listas de acceso, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. ¿Cuál es la importancia de las listas de acceso?
2. ¿Cuál es la mejor forma de identificar las listas de acceso en IPv6 con letras o con números?
Explique.

3. ¿Qué indica la siguiente línea de código? Justifique su respuesta

```
IPv6 access-list denegarpingftp
Deny icmp 2008::/64 2006::/64 echo-request
Permit icmp any any
Permit ipv6 any any
```

4. ¿Cuál es la diferencia entre las listas de acceso de IPv6 e IPv4?
5. Define ¿Para qué se utiliza el comando IPv6 traffic-filter?
6. Indique cual es la afirmación correcta en las siguientes enunciados
 - a) Las listas de acceso IPv6 no funcionan si activas listas de acceso de IPv4 sobre una determinada interfaz.
 - b) El identificador de las listas de acceso en IPv6 no pueden ser numéricos
 - c) No puedes tener más de 2 listas de acceso IPv6 activadas sobre una misma interfaz o subinterfaz.
 - d) Las listas de acceso en IPv6 no tienen una forma estándar ni extendida. El comando IPv6 access-list 1 define una lista de acceso que se creara con el identificador 1.
 - e) Ninguna de las anteriores.



PRÁCTICA 7: PROTOCOLO DE ENRUTAMIENTO INTERNO (RIPng, EIGRPv6, OSPFv3).

Objetivo general

- Implementar los protocolos de enrutamiento internos en IPv6.

Objetivos específicos

- Determinar y establecer el proceso de configuración de cada uno de los protocolos de enrutamiento interno.
- Definir e implementar el método de distribución de rutas en los dispositivos enrutadores, con la finalidad de comunicar y asociar los diferentes protocolos internos.
- Configurar DHCPv6 para la asignación dinámica de direcciones IPv6, así como la autoconfiguración.

Introducción

Los protocolos de enrutamientos dinámicos son un conjunto de reglas, por el cual los routers comparten su información. Cuando un router recibe los datos de nuevas rutas de acceso o la modificación de una de ellas, la tabla de enrutamiento que posee el dispositivo se actualiza de manera automática, permitiendo lograr la comunicación entre sí. Los protocolos de enrutamientos interno poseen características que los diferencia entre sí, por lo cual cuando un router necesita actualizar su información con un protocolo diferente al que está ejecutando, es requerido proporcionarles la redistribución de rutas con la finalidad de conectar los diferentes protocolos en el router.

Los protocolos de enrutamiento dinámico interno son los siguientes:

RIPng: Es un protocolo de enlace interno, mediante el cual los routers pertenecientes a un mismo Sistema Autónomo intercambian y actualizan su información acerca de las redes a las que se encuentran conectados.

OSPFv3: Es un protocolo de estado de enlace, en oposición al protocolo de vector de distancia, toma las decisiones de enrutamiento basado en los estados de los enlaces que conectan los equipos de origen y destino. El estado de un enlace es una descripción de esa interfaz y su relación con sus dispositivos de red vecinos. La información de interfaz incluye el prefijo IPv6 de la interfaz, la máscara de red, el tipo de red que está conectado y los dispositivos conectados a esa red. Esta información se propaga en varios tipos de anuncios de estado de enlace llamado LSA.

EIGRPv6: Es un protocolo de estado-enlace, desarrollado por Cisco, basado principalmente en el vector-distancia, obtiene la información de la redes a través de los routers vecinos conectado directamente a ellos,



por lo tanto cuando un router detecta un nuevo vecino, incluye su dirección IP y la interfaz por la que está conectado en la tabla de vecinos y mantiene la tabla de la topología con todos los destinos recibidos por este vecino.

La realización de este laboratorio le permitirá al estudiante comprender de manera correcta y eficiente como trabaja cada protocolo interno y de los requerimientos necesario para lograr correcta configuración.

Requerimientos:

Hardware	Software
➤ Computadora con los siguientes requisitos: • Procesador mínimo de velocidad de 2.1 GHz. • Memoria RAM de 1 GB.	➤ Simulador Packet Tracer 6.1.1 : • 12 Routers. • 13 Pc. • 8 switch.

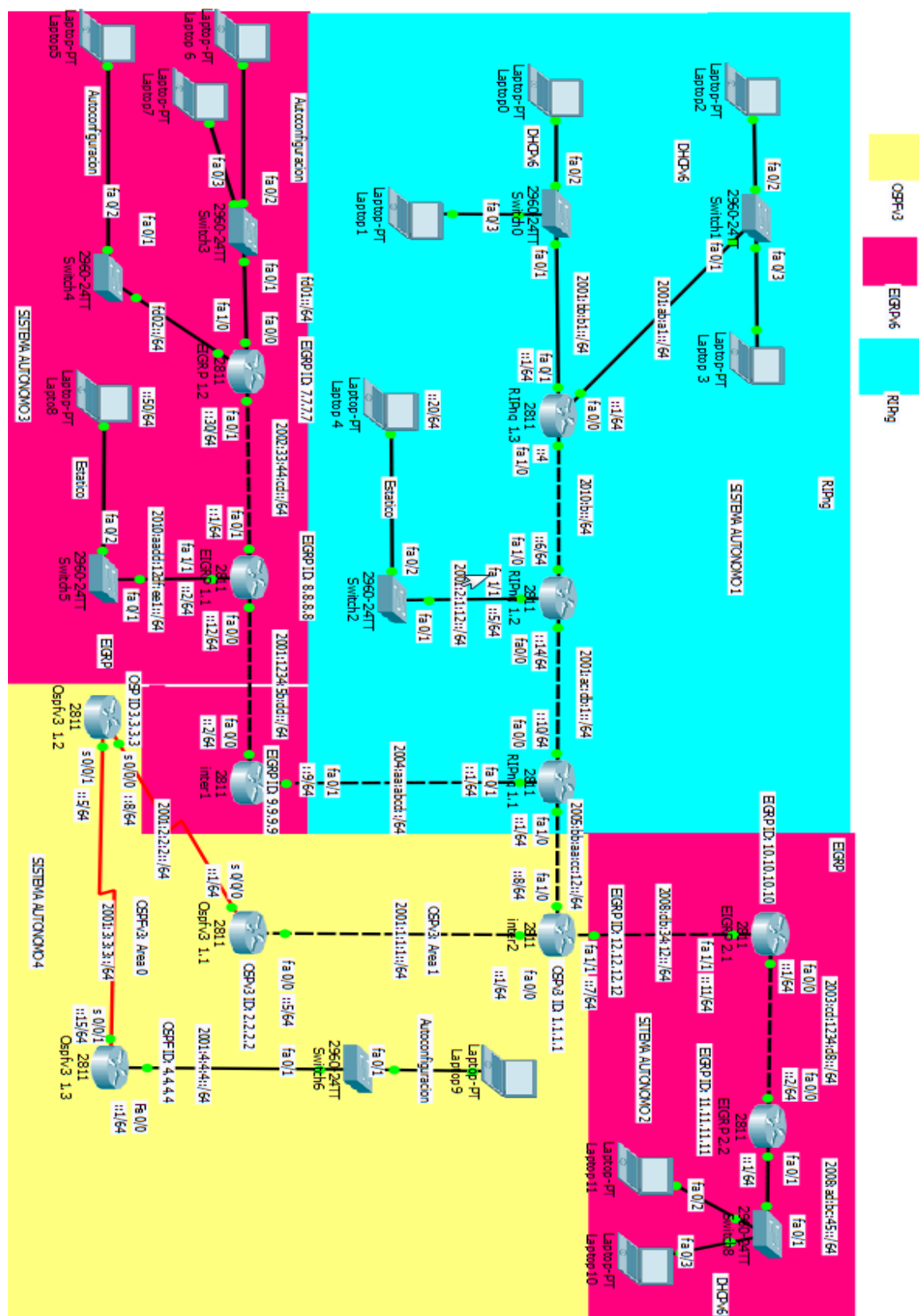
Conocimientos previos

Para la correcta realización de esta práctica el estudiante deberá tener conocimientos básicos:

- Asignación de direcciones IPv6.
- Configuración de enrutamiento estático IPv6.
- Configuración servidores DHCP para IPv6.



Topología





Funcionalidad

En la imagen anterior se observa una topología de red, en la cual se aplicaron los protocolos de enrutamiento interno que se pueden implementar en IPv6 mediante el simulador Packet Tracer.

La topología está estructurado de la siguiente manera:

Sistema Autónomo 1: Se implementó el protocolo de enrutamiento dinámico interno RIPng, los equipos finales se le asignó direcciones de manera estática y dinámica. Para la asignación de direcciones dinámica fue requerido el uso del DHCPv6.

Sistema Autónomo 2: Se implementó el protocolo de enrutamiento dinámico interno EIGRPv6, cada router posee un ID único. Los equipos finales se le asignó direcciones de manera estática y dinámica. Para la asignación de direcciones dinámica fue requerido el uso del DHCPv6.

Sistema Autónomo 3: Se implementó el protocolo de enrutamiento dinámico interno EIGRPv6, cada router posee un ID único. Los equipos finales se le asignó direcciones de manera estática y de autoconfiguración.

Sistema Autónomo 4: Se implementó el protocolo de enrutamiento dinámico interno OSPFv3, cada Router posee un ID único y todo pertenecen al área "1". Los equipos finales se le asignó direcciones de manera de autoconfiguración.

NOTA: En los sistemas Autónomos se configura la redistribución de rutas, específicamente en los siguientes routers:

- Routers => inter1: Redistribución de RIPng con EIGRP
- Routers => inter2: Redistribución de RIPng, EIGRP y OSPFv3.

Comandos generales de IPv6

Comando o Acción	Descripción
enable	Habilita el modo EXEC privilegiado.
configure terminal:	Entra en el modo de configuración global.
ipv6 unicast-routing	Permite el envío de datagramas IPv6 unicast.
interface type number	Configura un tipo de interfaz y entra en la interfaz el modo de configuración
ipv6 enable	Permite el procesamiento de IPv6 en una interfaz.
ipv6 address {ipv6-address/prefix-length prefix-name sub-bits/prefix-length}	Configura una dirección IPv6 en base a un prefijo y permite el procesamiento de IPv6 en una interfaz.
DHCPv6	



ipv6 dhcp pool pool-name Example: Router(config)# ipv6 dhcp pool pool1	Configura un conjunto de información de configuración DHCPv6 y entra en el modo de configuración de la agrupación DHCPv6.
prefix-delegation pool pool-name [lifetime valid-lifetime preferred-lifetime] Example: Router(config-dhcp)# prefix-delegation pool pool1 lifetime 1800 60	Especifica un prefijo IPv6 pool local llamado desde la que los prefijos se delegan a los clientes DHCPv6.
exit	Exits DHCPv6 pool configuration mode configuration mode, and returns the router to global configuration mode.
interface type number	Especifica un tipo de interfaz y el número, y coloca el router en el modo de configuración de interfaz.
interface type number	Especifica un tipo de interfaz y el número, y coloca el router en el modo de configuración de interfaz.
ipv6 dhcp server pool-name [rapid-commit] [preference value] [allow-hint] Example: Router(config-if)# ipv6 dhcp server pool1	Enables DHCPv6 on an interface.
ipv6 dhcp server pool-name [rapid-commit] [preference value] [allow-hint] Example: Router(config-if)# ipv6 dhcp server pool1	Enables DHCPv6 on an interface.
Router(config)#ipv6 local pool pool1 2001:4000::/40 64 Example: IPv6 local pool pool1 2001:4000::/40 64	Define un pool de bloque de direcciones que se le asignaran al cliente.
Protocolos de Enrutamiento	
RIPng	
ipv6 router rip Word	Habilita enrutamiento RIPng en el router
ipv6 rip Word enable	Activa el enrutamiento RIPng en una interfaz
OSPFv3	
ipv6 router ospf process-id	Habilita el enrutamiento OSPFv3 en el router
router-id A.B.C.D	Asigna un id al router. El ID debe de ser una dirección de tipo IPv4



ipv6 ospf process-id <i>área</i> área-id	Activa el enrutamiento OSPFv3 en la interfaz. Además de definir el área al que va a pertenecer.
EIGRPv6	
ipv6 router eigrp process-id	Activa enrutamiento EIGRPv6 en el router.
eigrp router-id A.B.C.D	Designa un id al router. Este ID tiene que ser una dirección IPv4.
no shutdown	Habilita el proceso EIGRPv6 que por defecto esta desactivado.
ipv6 eigrp process-id	Activa el enrutamiento EIGRPv6 en la interfaz.
Redistribución de Protocolos de Enrutamiento	
redistribute <i>protocolo</i> [<i>proceso-id</i>] [metric { <i>valor-metrica</i> transparent }] include-connected	Configura la redistribución en el protocolo especificado
redistribute connected	Permite anunciar por el tipo de protocolo las interfaces directamente conectados.

Direccionamiento IPv6 / Sistema Autónomo 1					
Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
RIPng 1.1	2001:AC:DB:1 ::10/64	2004:AA:ABCD: :1/64	2005:BB:AA:CC :12::1/64		RIPng
RIPng 1.2	2001:AC:DB:1 ::14/64		2010:B::6/64	2001:2:1:12::5/6 4	RIPng
RIPng1.3	2001:AB:A1::1 /64	2001:BB:B1::1/6 4	2010:B::4/64		RIPng
Inter1		2004:AA:ABCD: :9/64			RIPng
Inter2			2005:BB:AA:CC :12::8/64		RIPng

Direccionamiento IPv6 / Sistema Autónomo 2					
Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
EIGRP2.1	2003:CD:1234 :D8::1/64			2008:DB:34:12:: 11/64	EIGRPv6



EIGRP2.2	2003:CD:1234: :D8::2/64	2008:AD:BC:45: :1/64			EIGRPv6
Inter2				2008:DB:34:12:: 7/64	EIGRPv6

Direccionamiento IPv6 / Sistema Autónomo 3					
Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
EIGRP1.1	2001:1234:5B: DD::12/64	2002:33:44:CD:: 1/64		2010:AADD:12D F:EE1::2/64	EIGRPv6
EIGRP1.2	FD01::/64	2002:33:44:CD:: 30/64	FD02::/64		EIGRPv6
Inter1	2001:1234:5B: DD::2/64				EIGRPv6

Direccionamiento IPv6 / Sistema Autónomo 4						
OSPFv3						
Routers	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	S 0/0/0	S 0/0/1
Ospf3- 1.1	2001:1:1:1::5/64				2001:2:2:2: :1/64	
Ospf3_ 1.2					2001:2:2:2: :8/64	2001:3:3:3 ::5/64
Ospf3_ 1.3	2001:4:4:4::1/64					2001:3:3:3 ::15/64
POOLS						
Pool Name	Gateway	Dirección de Red	Subred	Inicio Dirección IP		
Administración	2001:AB:A1::1/64	2001:AB:A1::/40	24	2001:AB:A1::2/64		
Contabilidad	2001:BB:B1::1/64	2001:BB:B1::/40	24	2001:BB:B1::2/64		
Informática	2008:AD:BC:45::1/64	2008:AD:BC:45::/40	24	2008:AD:BC:45::2/64		



Servidores DHCP				
Routers	Interfaz	Subinterfaces	Dirección IP	Gateway
RIPng 1.3	F 0/0		2001:AB:A1::1/64	2001:AB:A1::1/64
	F 0/1		2001:BB:B1::1/64	2001:BB:B1::1/64
EIGRP 2.2	F 0/1		2008:AD:BC:45::1/64	2008:AD:BC:45::1/64

Dispositivo Estático				
Equipo	S. Autónomo	Dirección Red	Dirección IP	Gateway
Laptop 4	1	2001:2:1:12::/64	2001:2:1:12::20	2001:2:1:12::5/64
Lapto8	3	2010:AADD:12DF: :EE1::/64	2010:AADD:12DF:E E1::50	2010:AADD:12DF:EE 1::2

ENUNCIADO

Asignación de direcciones IP.

Se deberá asignar las direcciones IP a las interfaces de los routers tal y como aparecen en el cuadros llamados: **Direccionamiento IPv6 / Sistema Autónomo 1-4**.

Configuración de equipos estático.

Se deberá asignar las direcciones IP a las interfaces de PC escritorio y laptops, tal como aparecen en el cuadro llamado: **Dispositivo Estático**

Configuración DHCPv6 y Autoconfiguración.

Para la asignación de direcciones dinámicas se deberá crear los Pools tal y como se muestra en los siguientes cuadros: **Pools y Servidores DHCP**.

La autoconfiguración se realizarán en los siguientes routers: **EIGRP 1.2** y **OSPFv3_1.3**. Los equipos finales que están conectado a los routers ante mencionado se le asigna direcciones IP por autoconfiguración



Configuración de protocolo de enrutamiento dinámico interno

Se deberá asignar el protocolo definido del siguiente cuadro: **Direccionamiento IPv6 / Sistema Autónomo 1-4.**

Redistribución de rutas

La redistribución de rutas se realizara en los siguientes routers:

- Routers => inter1: Redistribución de RIPng con EIGRPv6 (ambos sentidos).
- Routers => inter2: Redistribución de RIPng, EIGRPv6 y OSPFv3 (ambos sentidos).

Para la correcta configuración de comandos, se puede observar el cuadro llamado: **Comando Generales de IPv6.**

Tiempo estimado de solución

- 8 horas

Preguntas de Análisis

En los siguientes apartados se pretende que los alumnos sean capaces de analizar y responder las siguientes preguntas con respecto al tema de protocolos de enrutamientos internos en IPv6, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. ¿Indique las diferencias entre los protocolos de encaminamiento IPv6 e IPv4?
2. ¿Cuál es la principal diferencia entre los protocolos dinámicos internos de IPv6 e IPv4 al configurarlos en un router?
3. ¿Por qué la declaración del comando router-id en los protocolos dinámicos internos de IPv6 exceptuando RIPng debe de ser única?
4. ¿Es igual el proceso de redistribución de rutas en IPv6 como era en IPv4? Explique.



5. De acuerdo a conocimientos adquiridos en estudios anteriores de IPv4 como estudios superiores de este tema IPv6 con referencia a los protocolos de enrutamiento dinámicos. Elige la afirmación correcta:

- a) La declaración de los protocolos dinámicos de enrutamiento en IPv6 es global e igual como en el direccionamiento IPv4
- b) Para activar los protocolos de enrutamiento IPv6 no es necesario el comando IPv6 unicast-routing
- c) El protocolo de enrutamiento dinámico IPv6 se activa directamente sobre las interfaces de un router y no pueden haber más de 1 encaminamiento dinámico IPv6.
- d) En las interfaces de un router pueden ejecutarse diferentes protocolos de encaminamiento IPv6 y en RIPng no es necesario el comando **router-id**.
- e) Ninguna de las anteriores.



PRÁCTICA 8: PROTOCOLOS DE ENRUTAMIENTO INTERNOS y EXTERNOS.

Objetivo general

- Configurar los protocolos de enrutamientos dinámicos en IPv6, y/o entender su interrelación entre cada uno de ellos.

Objetivos específicos

- Escribir la configuración de protocolos de enrutamiento internos (RIPng, ospfv3, eigrpv6 e is-isv6) dentro de diferentes sistemas autónomos.
- Establecer y configurar el protocolo externo bgpv4 en los enrutadores de borde de cada uno de los sistemas autónomos.
- Estudiar e implementar la redistribución de rutas en los routers de borde, para poder comunicar los diferentes sistemas autónomos.

Introducción

Al realizar esta práctica, se pretende que los alumnos sean capaces de entender y poder configurar los diferentes tipos de protocolos de enrutamiento dinámico internos (RIPng, OSPFv3, EIGRPv6 e IS-ISv6), y a su vez comunicar los diferentes sistemas autónomos a través del protocolo de enrutamiento externo BGPv4. Este protocolo se configurará en el/los routers de borde de cada Sistema Autónomo siendo estos, los routers que funcionarán con dos tipos de protocolo de enrutamiento: BGPv4 y el protocolo de enrutamiento interno que corresponda con el sistema autónomo.

Requerimientos:

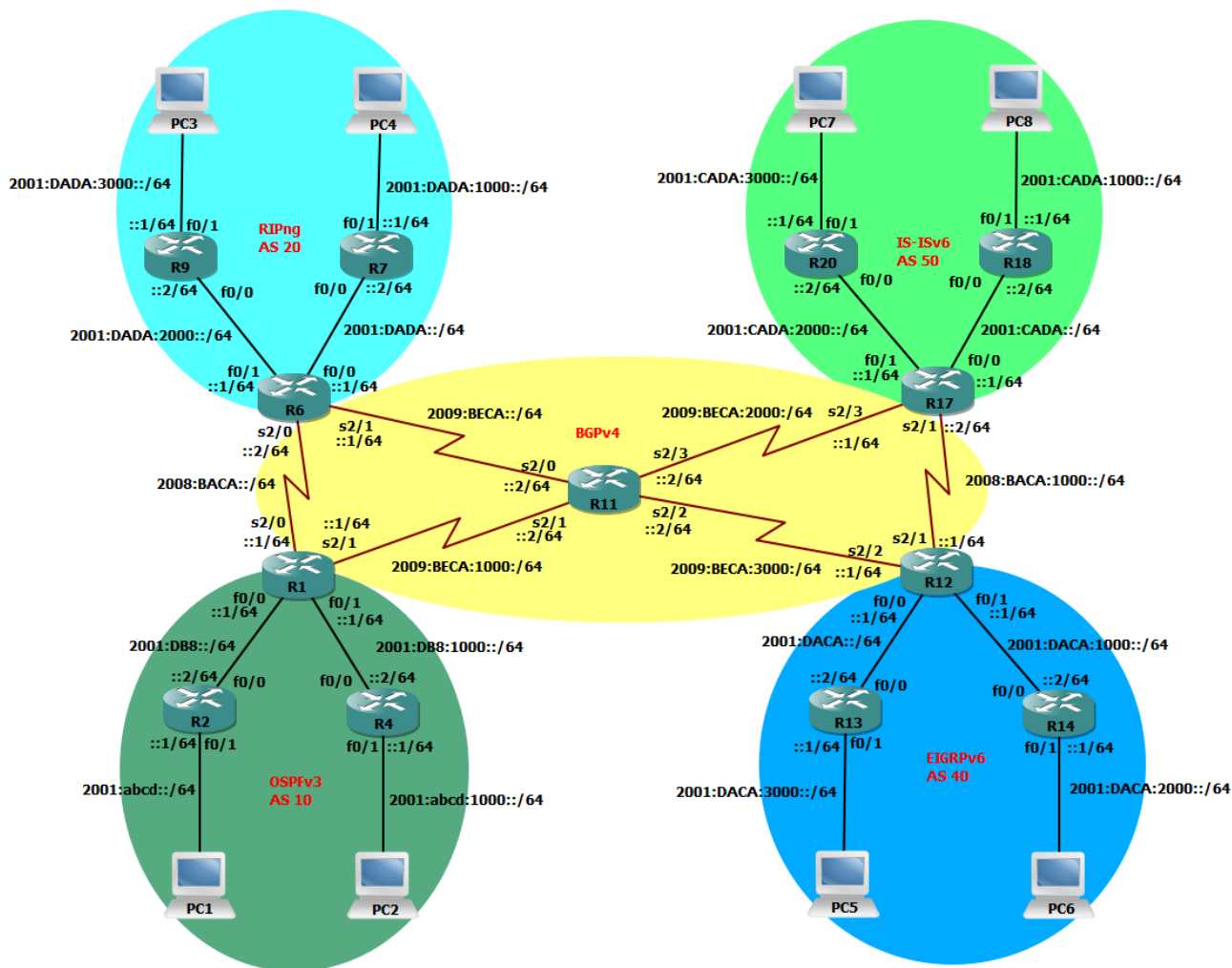
Hardware	Software
Computadora con los siguientes requisitos: ➤ Procesador mínimo de velocidad de 2.1 GHz ➤ Memoria RAM de 2 GB.	Emulador de redes GNS3 1.2.1 con los siguientes requerimientos: ➤ 10 Routers de la serie C7200 ➤ 3 Routers serie C3725 ➤ 8 PCs.

Conocimientos previos

Para la correcta realización e implementación de esta práctica es necesario que el estudiante tenga conocimientos básicos acerca del uso y funcionamiento de los diferentes tipos de protocolos de enrutamiento dinámico e implementación del protocolo IPv6 a su vez del protocolo de comunicación entre diferentes sistemas autónomos (BGPv4).



Topología



Funcionalidad

En la figura anterior se estarán aplicando diferentes tipos de protocolos de enrutamiento dinámicos combinándose con BGPv4 para su correcta comunicación entre cada uno de ellos.

- Se establecerá un protocolo de enrutamiento dinámico para cada Sistema Autónomo como se muestra en la figura anterior haciendo uso del protocolo IPv6.
- En los routers de borde de cada SA se implementa BGPv4 para la comunicación entre los distintos Sistemas Autónomos.
- Se hará uso de redistribución de rutas en los routers de borde entre BGPv4 y el que se designó en ese sistema autónomo.



NOTA: El router C3725 soporta EIGRP para IPv6 en cambio el C7200 no.

Comandos de ayuda

Comando	Descripción
RIPng	
ipv6 router rip word	Habilita enrutamiento RIPng en el router
ipv6 rip word enable	Activa el enrutamiento RIPng en una interfaz
OSPFv3	
ipv6 router ospf process-id	Habilita el enrutamiento OSPFv3 en el router
router-id A.B.C.D	Asigna un id al router. El ID debe de ser una dirección de tipo IPv4
ipv6 ospf process-id área área-id	Activa el enrutamiento OSPFv3 en la interfaz. Además de definir el área al que va a pertenecer.
EIGRPv6	
ipv6 router eigrp process-id	Activa enrutamiento EIGRPv6 en el router.
eigrp router-id A.B.C.D	Designa un id al router. Este ID tiene que ser una dirección IPv4.
no shut	Habilita el proceso EIGRPv6 que por defecto esta desactivado.
ipv6 eigrp process-id	Activa el enrutamiento EIGRPv6 en la interfaz.
IS-ISv6	
router isis	Entra en modo configuración del router para ISIS
net dirección-net	Configura título entidad de red(NET) para IS-IS
address family ipv6	Permite configurar IS-IS con familia de direcciones IPv6.
ipv6 router is-is word	Activa IS-IS con IPv6 en una interfaz.
BGPv4	
router bgp As-Number	Accede al modo configuración del router para BGP.
bgp router-id A.B.C.D	Asigna un id al router BGP. La ID es una dirección IPv4
ddress-family ipv6	Permite configurar BGP con familia de direcciones IPv6.
Neighbor dirección-ip next-hop-self	Define al router vecino como siguiente salto.
Neighbor dirección-ip remote-as AS-Number	Establece una relación de vecino BGP.
Neighbor dirección-ip activate	Activa la comunicación entre distintos routers BGP vecinos.



No synchronization	Deshabilita la publicación de exclusivamente de redes con las cuales están en condiciones de comunicarse usando una ruta aprendida por un protocolo de enrutamiento interior.
---------------------------	---

Comandos generales	
Comando	Descripción
IPv6 address <IPv6-prefix/prefix-length>	Asigna una dirección IPv6 a una interfaz.
Interface tipo número	Cambia de modo configuración global a una al modo de configuración interfaz
ipv6 unicast-routing	Permite configurar IPv6
No shutdown	Habilita una interfaz
redistribute <i>protocolo</i> [<i>proceso-id</i>] [metric { <i>valor-metrica</i> transparent }] include-connected	Configura la redistribución en el protocolo Especificado
redistribute connected	Permite anunciar por el tipo de protocolo las interfaces directamente conectados que forman parte de la red del tipo de protocolo
Show ipv6 route	Muestra la tabla de enrutamiento IP
show running-config	Muestra el archivo de configuración activo

Datos de los Dispositivos

Routers							
Nombres	Direcciones IPv6						
	Fa 0/0	Fa 0/1	Se 2/0	Se 2/1	Se 2/2	Se 2/3	Protocolos y A.S
R1	2001:db8::1/64	2001:db8:1000::1/64	2008:baca::1/64	2009:beca:1000::1/64			OSPFv3 BGPv4(A.S 10)
R2	2001:db8::2/64	2001:abcd::1/64					OSPFv3
R4	2001:db8:1000::2/64	2001:abcd:1000::1/64					OSPFv3
R6	2001:da-da::1/64	2001:dada:2000::1/64	2008:baca::2/64	2009:beca::1/64			RIPng BGPv4 (A.S 20)
R7	2001:da-da::2/64	2001:dada:1000::1/64					RIPng



R9	2001:da da:2000 ::2/64	2001:dada:3 000::1/64					RIPng
R11			2009:beca: :2/64	2009:beca:100 0::2/64	2009:b eca:30 00::2/6 4	2009:b eca:20 00::2/6 4	BGPv4 (A.S 30)
R12	2001:da ca::1/64	2001:daca:1 000::1/64		2008:baca:100 0::1/64	2009:b eca:30 00::1/6 4		EIGRPv6 BGPv4(A.S 40)
R13	2001:da ca::2/64	2001:daca:3 000::1/64					EIGRPv6
R14	2001:da ca:1000 ::2/64	2001:daca:2 000::1/64					EIGRPv6
R17	2001:ca da::1/64	2001:cada:2 000::1/64		2008:baca:100 0::2/64		2009:b eca:20 00::1/6 4	IS-ISv6 BGPv6(A.S 50)
R18	2001:ca da::2/64	2001:cada:1 000::1/64					IS-ISv6
R20	2001:ca da:2000 ::2/64	2001:cada:3 000::1/64					IS-ISv6

Interfaz FastEthernet				
Nombres	Direcciones IPv6	Puerto local	Puerto remoto	Dirección Host remoto
R2	2001:abcd::1/64			
R4	2001:abcd:1000::1/64			
R7	2001:dada:1000::1/64			
R9	2001:dada:3000::1/64			
R13	2001:daca:3000::1/64			
R14	2001:daca:2000::1/64			
R18	2001:cada:1000::1/64			
R20	2001:cada:3000::1/64			
PC1	2001:abcd::10/64	20503	10009	127.0.0.1
PC2	2001:abcd:1000::10/64	20504	10016	127.0.0.1
PC3	2001:dada:3000::10/64	20506	10026	127.0.0.1
PC4	2001:dada:1000::10/64	20505	10023	127.0.0.1
PC5	2001:daca:3000::10/64	20501	10003	127.0.0.1



PC6	2001:daca:2000::10/64	20502	10006	127.0.0.1
PC7	2001:cada:3000::10/64	20508	10036	127.0.0.1
PC8	2001:cada:1000::10/64	20507	10033	127.0.0.1

ENUNCIADO

Asignación de direcciones IPv6

Se asignará las direcciones de las PCs y a las interfaces FastEthernet de los routers así como aparece en el cuadro **Interfaz FastEthernet**.

NOTA: configurar los puertos locales y remotos de cada una de las PCs recuerde que esta configuración de puertos locales y puertos remotos se hacen con respecto a los puertos locales y remotos de cada una de las PCs que te brinda Virtual PC. Ver el cuadro llamado **interfaz FastEthernet** para configuración de puerto local y puerto remoto de las PCs en GNS3.

- | | | | |
|------|-------|--------|--------|
| ➤ R2 | ➤ R13 | ➤ VPC1 | ➤ VPC5 |
| ➤ R4 | ➤ R14 | ➤ VPC2 | ➤ VPC6 |
| ➤ R7 | ➤ R18 | ➤ VPC3 | ➤ VPC7 |
| ➤ R9 | ➤ R20 | ➤ VPC4 | ➤ VPC8 |

De igual forma se le estarán asignando las direcciones IPv6 tanto a las interfaces seriales como FastEthernet a los routers tal como se muestra en la tabla **Routers**.

- | | | | |
|------|-------|-------|-------|
| ➤ R1 | ➤ R7 | ➤ R13 | ➤ R20 |
| ➤ R2 | ➤ R9 | ➤ R14 | |
| ➤ R4 | ➤ R11 | ➤ R17 | |
| ➤ R6 | ➤ R12 | ➤ R18 | |

Asignación de protocolo de enrutamiento (OSPFv3)

Se deberá de configurar este protocolo para enrutar el tráfico entre cada una de las redes existente del Sistema Autónomo 10 .Para lo cual vaya al cuadro **Routers** y observe los datos para los routers (**R1, R2, y R4**) pertenecientes a este Sistema Autónomo y configurar el Protocolo antes mencionado.

- | | | |
|------|------|------|
| ➤ R1 | ➤ R2 | ➤ R4 |
|------|------|------|



Asignación de protocolo de enrutamiento (RIPng)

Para comunicar las redes existentes del Sistema Autónomo 20, se estará usando RIPng como protocolo de enrutamiento, en cada uno de los routers (**R6, R7 y R9**) pertenecientes a este Sistema Autónomo. Observe tabla **Routers**.

➤ **R6**➤ **R7**➤ **R9**

Asignación de protocolo de enrutamiento (EIGRPv6)

Se debe configurar este protocolo para enrutar el tráfico entre cada una de las redes existente del Sistema Autónomo 40 vaya al cuadro de **Routers** para ver cuáles son los routers pertenecientes a este Sistema Autónomo y configurar el protocolo antes mencionado.

➤ **R12**➤ **R13**➤ **R14**

Asignación de protocolo de enrutamiento (IS-IS v6)

Se configurara el protocolo ISISv6 en el Sistema Autónomo 50 para enrutar el tráfico perteneciente a ese Sistema Autónomo ver cuadro de **Routers** para ver cuáles son los routers pertenecientes a este Sistema Autónomo y configurar el protocolo antes mencionado.

➤ **R17**➤ **R18**➤ **R20**

Configuración de Redistribución de rutas en los sistemas autónomos

La redistribución de rutas para que dos o más dispositivos (routers o switches capa 3) intercambien información de enrutamiento es preciso, que ambos equipos utilicen el mismo protocolo, sea RIP, EIGRP, OSPF, BGP, etc. Diferentes protocolos de enrutamiento, o protocolos configurados de diferente forma (por ejemplo. Diferente sistema autónomo en EIGRP) no intercambian información.

Por tanto, cuando un dispositivo aprende información de enrutamiento a partir de diferentes fuentes (protocolos), Cisco IOS permite que la información aprendida por una fuente sea publicada hacia otros dispositivos utilizando un protocolo diferente a esto es lo que se denomina "Redistribución" de rutas. Utilizar un protocolo de enrutamiento para publicar rutas que son aprendidas a través de otro medio.

NOTA: El sistema Autónomo 30 (A.S 30) solo estará configurado con BGPv4.El A.S 30 es el router R11 como se observa en la topología. Por lo tanto no se usa la redistribución de rutas al tener solo un protocolo de direccionamiento IPv6.



Redistribución de rutas en el sistema Autónomo 10 (A.S 10)

En el Sistema Autónomo 10 se deberá configurar redistribución de rutas en los routers de bordes para comunicar los equipos con las redes establecidas en este mismo con los demás Sistemas Autónomos.

NOTA: Recuerde que los routers de bordes de cada Sistema Autónomo son los únicos que se establecerán con dos o más tipos de protocolo de enrutamiento IPv6 (BGPv4 para comunicación con otros sistemas autónomos y el/los protocolo de enrutamiento interno del sistema autónomo).

Vaya a la tabla llamada **Routers** y la topología para observar cuales son los routers de bordes pertenecientes a este Sistema Autónomo y ver cuál/cuales es nuestro router vecino del Sistema Autónomo conectado directamente a nuestro Sistema Autónomo, después asignar la redistribución de rutas de BGPv4 en el protocolo OSPFv3 protocolo interno del sistema autónomo y de OSPFv3 en BGPv4 (*ver cuadro de ayuda de comandos para la correcta configuración*).

➤ R1

Redistribución de rutas en el sistema Autónomo 20 (A.S 20)

En el sistema autónomo 20 se configurará redistribución de rutas en los routers de bordes para comunicar los dispositivos con las redes establecidas en este sistema autónomo con otros sistemas autónomos.

Ver cuadro llamado **Routers** y topología para ver cuáles son los routers pertenecientes a este sistema autónomo ya que a diferencia de los demás sistema autónomos los dos únicos routers de este sistema autónomo son routers de bordes y se encuentran conectados directamente entre ellos por lo cual ellos son vecinos BGPv4 y se deberá configurar como tal, luego configurar la redistribución de rutas de BGPv4 en el protocolo RIPng protocolo interno del sistema autónomo y de RIPng en BGPv4.

➤ R6

Redistribución de rutas en el sistema Autónomo 40 (A.S 40)

En el sistema autónomo 40 se deberá configurar redistribución de rutas en los routers de bordes para comunicar los dispositivos con las redes establecidas en este sistema autónomo con los demás sistemas autónomos.

Vaya a la tabla llamada **Routers** y la topología para observar cuales son los routers de bordes pertenecientes a este sistema autónomo y ver cuál/cuales es nuestro router vecino del sistema autónomo conectado directamente a nuestro sistema autónomo, después asignar la redistribución de rutas de BGPv4 en el protocolo EIGRPv6 protocolo interno del sistema autónomo y de EIGRPv6 en BGPv4.

➤ R12

**Redistribución de rutas en el sistema autónomo 50 (A.S 50)**

En el sistema autónomo 50 se deberá configurar redistribución de rutas en los Routers de bordes para comunicar los dispositivos con las redes establecidas en este sistema autónomo con los demás sistemas autónomos.

Vaya a la tabla llamada **Routers** y la topología para observar cuales son los routers de bordes pertenecientes a este sistema autónomo y ver cuál/cuales es nuestro router vecino del sistema autónomo conectado directamente a nuestro sistema autónomo, después asignar la redistribución de rutas de BGPv4 en el protocolo IS-ISv6 protocolo interno del sistema autónomo y de IS-ISv6 en BGPv4.

➤ **R17**

Tiempo estimado de solución

➤ 8 horas

Preguntas de análisis

En los siguientes apartados se pretende que los alumnos sean capaces de analizar y responder las siguientes cuestiones en base al tema de protocolos de enrutamiento dinámico IPv6, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. Diga ¿Cuál es la diferencia en las siguientes líneas de códigos basadas en la redistribución de rutas BGPv4 en ISISv6? Justifique su respuesta.

Router isis 1

.....

!

Redistribute bgp 20 metric 10

!

Address-family ipv6

Redistribute bgp 20 metric 10

!

Router isis 1

.....

!



2. ¿Explique el uso de la línea de código siguiente en BGPv4 referente al uso de vecinos BGP?

```
Neighbor 2010:ACDC::1 activate
```

3. En OSPFv3. ¿Qué pasa si tenemos declarados 2 routers vecinos con el mismo identificador router-id?

¿Explique?

4. ¿Cuál de las siguientes afirmaciones que describen al enrutamiento de protocolo en IPv6 es

correcta?

- a) En un router pueden existir varios procesos ISISv6 y se declaran directamente sobre la/las interfaces.
- b) En OSPFv3 pueden existir diferentes procesos sobre las interfaces de un router y debe de tener varios identificadores para conectarse a un router vecino.
- c) Las redistribuciones de rutas no siempre se declaran en los routers vecinos para que se puedan comunicar diferentes tipos de protocolos de enrutamiento en IPv6.
- d) El comando **net** se usa como un identificador en el protocolo ISIS en ipv6.
- e) Ninguna de las anteriores.

5. ¿Por qué es necesario el comando **no shutdown** para declarar un proceso EIGRPv6?

6. El protocolo RIPng su identificar pueden ser números o letras. ¿Explique cuál es la mejor opción para declarar dicho protocolo?



PRÁCTICA 9: REDES VIRTUALES PRIVADAS.

Objetivo general

- Analizar y configurar una VPN (Virtual Private Network), para acceder a una red privada desde una infraestructura pública.

Objetivos específicos

- Configurar los parámetros de un túnel VPN para la comunicación segura entre dos routers.
- Implementar IPsec para el cifrado de datos y la autenticación de estos.

Introducción

En la siguiente práctica estudiaremos la funcionalidad de VPN sitio a sitio con IPsec mediante la interfaz de túnel virtual, utilizando cifrado de datos para proteger el tráfico IPv6 entre las dos redes de confianza. En esta práctica el estudiante podrá visualizar cada detalle de la configuración para su mejor entendimiento.

Requerimientos:

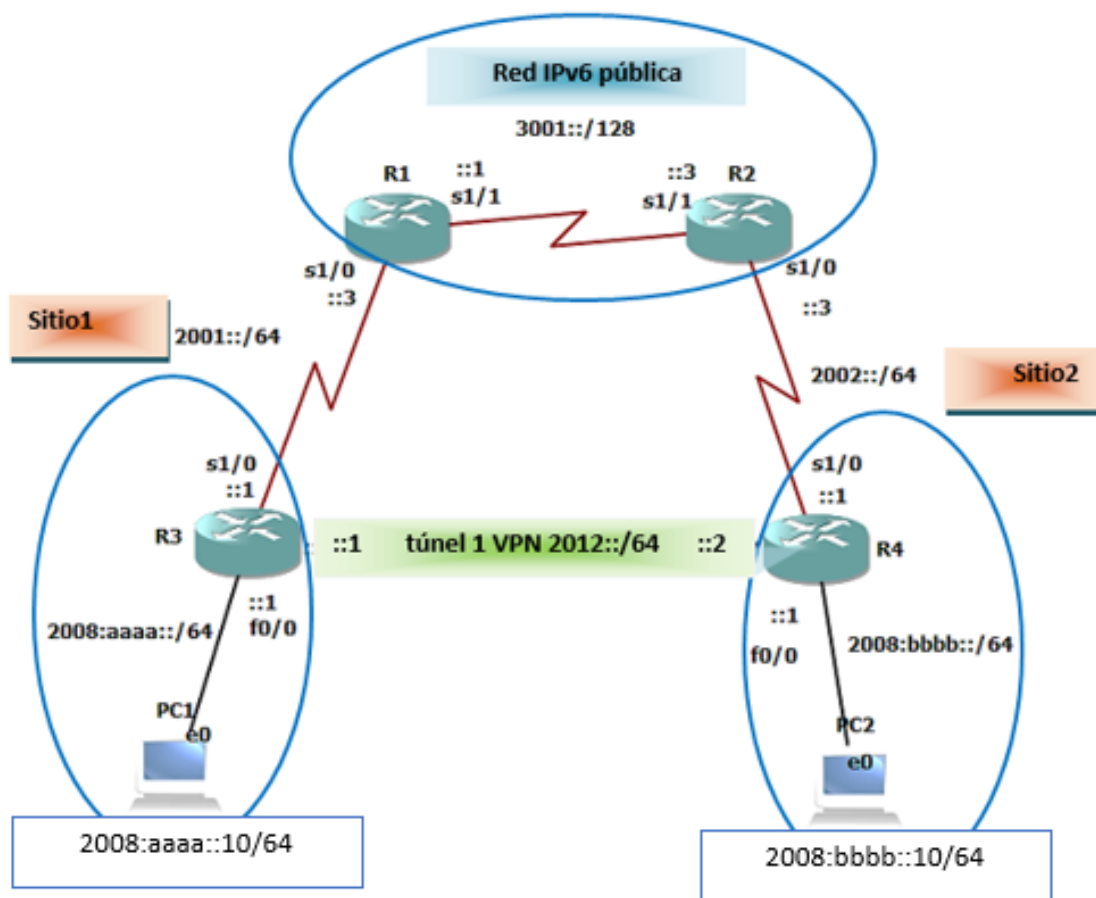
Hardware	Software
Computadora con los siguientes requisitos: ➤ Procesador con velocidad de 2.1 ghz ➤ Memoria RAM de 2 GB.	Emulador de redes GNS3 con los siguientes elementos: ➤ 4 Router c7200 ➤ 2 PCs

Conocimientos previos

Para la correcta realización de esta práctica el estudiante deberá tener conocimientos básicos tanto de los conceptos previos como de la configuración de IPsec al igual que VPN.



Topología



Funcionamiento

En la topología de esta práctica se establecerá lo siguiente:

Se deberá configurar protocolo de enrutamiento Rip para establecer comunicación entre los routers exteriores a la red.

Los routers principales del sitio1 y sitio2 son R3 y R4, estos router se comunican a través de una red pública externa, la cual no debe tener información de los datos que se envían dichos routers y se configuran los siguientes:

- Configurar la política IKE
- Configurar la política IPsec
- La clave de intercambio es: Telemática.
- Crear un perfil de ISAKMP.

**Comandos de ayuda**

Comandos	Definición
crypto isakmp policy <i>priority</i>	Habilitar política de seguridad
encryption {des 3des}	Especifica el algoritmo de encriptación que se utiliza en IKE.
group {1 2}	Especifica el identificador de grupo del algoritmo de Diffie Hellman.
authentication {rsa-sig rsa-encr pre-share}	Método de autenticación para el intercambio de claves.
crypto isakmp key <i>clave</i> address <i>peer-address</i>	Comparte la clave con el par especificado.
crypto ipsec transform-set transform-set-name transform1 [transform2 [transform3]]	Define el protocolo de transform-set, selecciona si se utiliza AH o ESP.
set transform-set transform-set-name	De los transform-set que se han definido especifica cual aplica al túnel.
Show crypto isakmp sa	Muestra las sesiones activas en el router.
show crypto engine connection active	Muestra la información de configuración de la criptografía.
Show interfaces tunnel id-tunnel	Muestra la información de interfaz especificada.

Datos de los dispositivos

Red publica			
Dispositivos	Interfaz	Dirección IP	Dirección de Red
R1	S1/0	2001::3	2001::/64
	S1/1	3001::1	3001::/128
R2	S1/0	2002::3	2002::3/64
	S1/1	3001::3	3001::/128



Redes de sitio1y sitio2			
Dispositivos	Interfaz	Dirección IP	Dirección de Red
R3	S1/0	2001::1	2001::/64
	F0/0	2008:aaaa::1	2008:aaaa::/64
PC1	e0	2008:aaaa::10	2008:aaaa::/64
R4	S1/0	2002::1	2002::/64
	F0/0	2008:bbbb::1	2008:bbbb::/64
PC2	e0	2008:bbbb::10	2008:bbbb::/64

Túnel VPN			
Dispositivos	Interfaz	Dirección IP	Dirección de Red
R3	S1/0	2012::1/64	2012::/64
R4	S1/0	2012::2/64	2012::/64

ENUNCIADO

Crear la topología en el Emulador GNS3, asignándole las direcciones ip establecidas en las tablas correspondiente a cada dispositivo.

Configuración de la VPN en los routers R3 y R4.

1. Configuración del túnel VPN

- Habilitar un túnel VPN para el R3 para que puedan acceder desde una infraestructura pública a R4.
- Asignar la dirección del túnel mostradas en la tabla **Túnel VPN**.
- Habilitar el origen y destino en cada extremo.
- Configurar el túnel modo IPSec ipv6.

2. Configurar la política de IKE.

- Definir la política de encriptación IKE con prioridad 10.
- El algoritmo de encriptación que se utiliza es 3des.
- El identificador del grupo del algoritmo de Diffie Hellman definir el 2.
- El método de autenticación a utilizar es pre-shared.



3. Configurar la clave de intercambio.

En este proceso los routers R3 y R4 deben configurarse con la misma clave a utilizar para el intercambio entre ambos extremos, la clave definida es **TELEMATICA**.

4. Configuración del protocolo de transform-set.

- En este proceso se eligió el protocolo de encriptación esp, el algoritmo 3des y como algoritmo de autenticación esp-sha.hmac, el nombre del transform-set es **TLMCA**.

NOTA: Es importante que el nombre que se le asignó al protocolo coincida con el que se configura en cada extremo del túnel para que haya conectividad.

5. Asignar las rutas de forma estáticas en R3 y R4.

6. Verificar los detalles de la configuración.

- Visualizar las sesiones ISAKMP en cada router y la información de la criptografía.
- Desde los routers R3 y R4 realizar un ping especificando la dirección origen del router y la dirección de destino.
- Luego verificar con el comando traceroute.

Tiempo estimado de solución

- 6 horas.

Preguntas de análisis

1. ¿Por qué es necesario implementar IPsec en una VPN?
2. ¿Cómo se pueden visualizar las asociaciones de seguridad en cada router y la información de la criptografía?
 - a. #show crypto isakmp sa
 - b. #show crypto ipsec sa
 - c. #show crypto engine connections active



3. ¿Con que comando se creó la política de encriptación?
 - a. `#crypto isakmp policy priority`
 - b. `#crypto isakmp enable`
 - c. `#crypto map map-name seq-num ipsec-isakmp`

4. Al configurar la clave de intercambio
 - a. Se debe configurar la misma clave en ambos extremos.
 - b. La clave puede diferente en cada extremo
 - c. Se configura en un solo extremo

5. ¿Cuál de los siguientes métodos de autenticación se utilizó?
 - d. `rsa-sig`
 - e. `rsa-encr`
 - f. `pre-share`



PRÁCTICA 10: VOIP CON IPv6

Objetivo general:

- Implementar VOIP en redes IPv6 aplicando Nat64

Objetivos específicos:

- Determinar las condiciones que permitan a Nat64 traducir direcciones en ambos sentidos (IPv4 e IPv6).
- Implementar VOIP en redes IPv6 usando como traductor de direcciones IP Nat64.

Introducción:

VOIP es un sistema de telefonía que ofrece llamadas de voz a través de redes de datos, actualmente basa su funcionamiento a través de IPv4, debido a la evolución del internet, que trajo consigo la implementación de IPv6, este sistema se está reestructurando de manera que trabaje eficientemente con este nuevo protocolo, por lo que por momento es fiable utilizar Nat64 como traductor de direcciones que le permite a VOIP actuar de manera eficiente sobre redes IPv6.

Requerimientos:

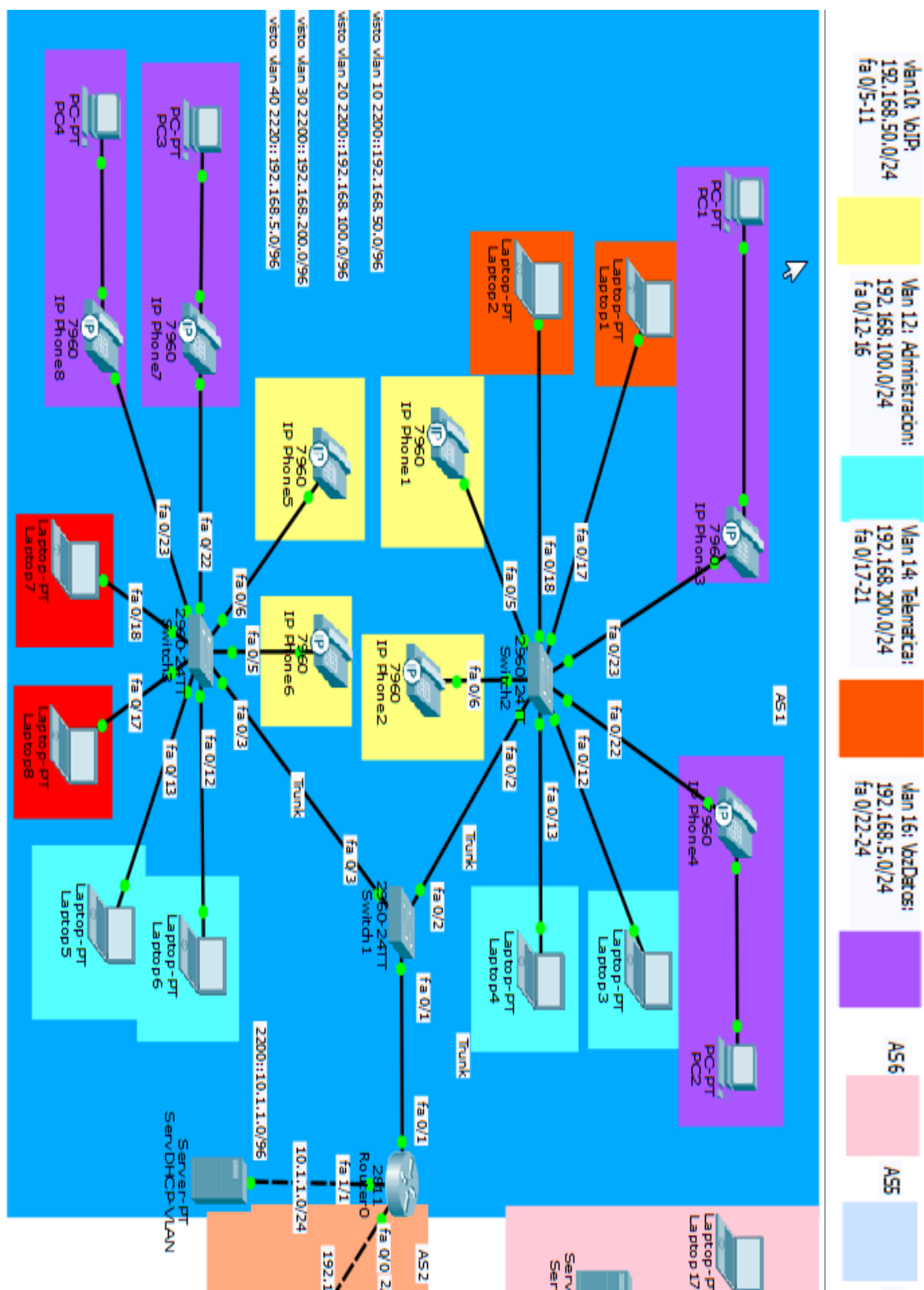
Hardware	Software
Computadora con los siguientes requisitos: ➤ Procesador mínimo de velocidad de 2.1 GHz ➤ Memoria RAM de 1 GB.	Simulador Packet Tracer 6.1.1 con los siguientes requerimientos: ➤ 10 Routers ➤ 19 laptops ➤ 4 Pc ➤ 3 Servidores ➤ 7 Switch.

Conocimientos previos:

Para la correcta realización de esta práctica el alumno debe tener conocimientos de direccionamiento IPv6, enrutamiento estático y dinámico IPv6, DHCPv6, direccionamiento IPv4, enrutamiento estático y dinámico de IPv4.



Topología:





**Funcionalidad:**

En la figura anterior se estableció VOIP dentro de una estructura IPv4 con la finalidad de aplicar un mecanismo de transición como Nat64 para interactuar dentro de una red IPv6.

La topología está estructurado de la siguiente manera:

- **Sistema Autónomo 1:** Se implementó una estructura de red IPv4, dentro de los cual se crearon múltiples VLANs, se creó un servidor DHCP que contienen rango de direcciones de cada una de las VLANs y el protocolo de enrutamiento aplicado fue RIPv2.
- **Sistema Autónomo 2:** Se implementó el Nat64 para lograr la comunicación entre IPv6 e IPv4.
- **Sistema Autónomo 3:** Se implementó una estructura de red IPv6 con el protocolo RIPv6.
- **Sistema Autónomo 4:** Se implementó una estructura de red IPv6 con el protocolo EIGRPv6.
- **Sistema Autónomo 5:** Se implementó una estructura de red IPv6 con el protocolo OSPFv3, además se crearon VLANs.
- **Sistema Autónomo 6:** Se implementó una estructura de red IPv6 con el protocolo RIPv6.

Además:

- Configurar VTP en los switches perteneciente a la redes IPv4 y VLANs estática en redes IPv6.
- aplicar DHCP (IPv6 e IPv4) en las redes configuradas.

Comandos de ayuda

Comando o Acción	Descripción
enable	Habilita el modo EXEC privilegiado.
configure terminal:	Entra en el modo de configuración global.
ipv6 unicast-routing	Permite el envío de datagramas IPv6 unicast.
interface type number	Configura un tipo de interfaz y entra en la interfaz el modo de configuración
ipv6 enable	Permite el procesamiento de IPv6 en una interfaz.
ipv6 address {ipv6-address/prefix-length prefix-name sub-bits/prefix-length}	Configura una dirección IPv6 en base a un prefijo y permite el procesamiento de IPv6 en una interfaz.
DHCPv6	



ipv6 dhcp pool pool-name Example: Router(config)# ipv6 dhcp pool pool1	Configura un conjunto de información de configuración DHCPv6 y entra en el modo de configuración de la agrupación DHCPv6.
prefix-delegation pool pool-name [lifetime valid-lifetime preferred-lifetime] Example: Router(config-dhcp)# prefix-delegation pool pool1 lifetime 1800 60	Especifica un prefijo IPv6 pool local llamado desde la que los prefijos se delegan a los clientes DHCPv6.
Exit	Sale del modo configuración de DHCPv6, y regresa al modo configuración global del router.
interface type number	Especifica un tipo de interfaz y el número, y coloca el router en el modo de configuración de interfaz.
ipv6 dhcp server pool-name [rapid-commit] [preference value] [allow-hint] Example: Router(config-if)# ipv6 dhcp server pool1	Activa DHCPv6 sobre una interfaz.
Router(config)#ipv6 local pool pool1 2001:4000::/40 64 Example: IPv6 local pool pool1 2001:4000::/40 64	Define un pool de bloque de direcciones que se le asignaran al cliente.
Protocolos de Enrutamiento	
RIPng	
ipv6 router rip Word	Habilita enrutamiento RIPng en el router
ipv6 rip Word enable	Activa el enrutamiento RIPng en una interfaz
OSPFv3	
ipv6 router ospf process-id	Habilita el enrutamiento OSPFv3 en el routers
router-id A.B.C.D	Asigna un id al Routers. El ID debe de ser una dirección de tipo IPv4
ipv6 ospf process-id área área-id	Activa el enrutamiento OSPFv3 en la interfaz. Además de definir el área al que va a pertenecer.
EIGRPv6	
ipv6 router eigrp process-id	Activa enrutamiento EIGRPv6 en el router.
eigrp router-id A.B.C.D	Designa un id al router. Este ID tiene que ser una dirección IPv4.
no shutdown	Habilita el proceso EIGRPv6 que por defecto esta desactivado.
ipv6 eigrp process-id	Activa el enrutamiento EIGRPv6 en la interfaz.



Redistribución de Protocolos de Enrutamiento	
redistribute <i>protocolo</i> [<i>proceso-id</i>] [metric { <i>valor-metrica</i> transparent }] include-connected	Configura la redistribución en el protocolo especificado
redistribute connected	Permite anunciar por el tipo de protocolo las interfaces directamente conectados.
Configuración VLANS	
vlan database	Accede a la base de datos de VLANs
vlan [<i>número-vlan name nombre-vlan</i>]	Crea una VLAN especificando el nombre y el número.
encapsulation dot1q <i>vlan-id</i>	Establece el método de encapsulación de la interfaz de enlace troncal 802.1Q VLAN, también especifica el ID de VLAN para la que las tramas deben ser etiquetados
switchport access <i>vlan vlan-id</i>	Asigna la VLAN por defecto para un puerto
option 150 ip gateway	Campo obligatorio para implementar correctamente el servicio de VOIP
show vlan	Muestra información de las VLANs

Nat64	
ipv6 nat v4v6 source <i>ipv4-address ipv6-address</i>	Traduce la dirección IPv4 a una dirección Ipv6
ipv6 nat v6v4 source <i>ipv6-address ipv4-address</i>	Traduce la dirección IPv6 a una dirección Ipv4
ipv6 nat prefix <i>ipv6-prefix/prefix-length</i> R1(config)# ipv6 nat prefix 1144::/96	Especifica que los paquetes que coincidan con esa dirección sean traducidos. Es importante señalar que la longitud de prefijo debe ser 96 para señalar redes IPv4.

Datos de los equipos

Direccionamiento IPv4 / Sistema Autónomo 1					
Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
Router0	192.168.1.2/30	Subinterfaces: 0/1.10:		10.1.1.1/24	RIPv2



		192.168.50.1/24 0/1.12: 192.168.100.1/24 0/1.14: 192.168.200.1/24 0/1.16: 192.168.5.1/24			
--	--	---	--	--	--

Direccionamiento IPv6 / Sistema Autónomo 2				
Equipo	Fa 0/0	Fa 0/1	Loopback0	Enrutamiento
Router1 NAT64	192.168.1.1/3 0	3010::1/64	12.12.12.1/24	RIPv2 RIPng

Direccionamiento IPv6 / Sistema Autónomo 3					
Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
Router2	3001:DB8:B::2 0/64	3010::3/64	2017:11:22::1/6 4	2018::50/64	Fa 0/0: EIGRPv6 Fa 0/1: RIPng Fa 1/0 : RIPng Fa 1/1: RIPng
	interface Loopback0		12.12.12.3/24		RIPv2
Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
Router3				2018::60/64	RIPng
Router4	2010:DB:1:1:: 2/64		2017:11:22::2/6 4		RIPng
Router5	2010:DB:1:1:: 1/64	2001:AB:2:2::1/ 64	2020:DB:AB:88: :1/64		RIPng



Direccionamiento IPv6 / Sistema Autónomo 4					
Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
Router6	3001:DB8:B::1 /64	2020:1234:22:1 1::90/64		2005::5/64	EIGRPv6
Router7		2007::2/64		2005::8/64	Fa 0/1: RIPng Fa 1/1: EIGRPv6
Router8		2007::1/64			RIPng

Direccionamiento IPv6 / Sistema Autónomo 5					
Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
Router9	Subinterfaces: 0/0.4: 3007:1:1::1/64 0/0.6: 3007:2:2::1/64	Subinterfaces: 0/1.8: 3006:1:1::1/64 0/1.12: 3006:2:2::1/64	3002:33:44:B::5/ 64		OSPFv3
Router8			3002:33:44:B::4/ 64		OSPFv3

Direccionamiento IPv6 / Sistema Autónomo 6					
Equipo	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1	Enrutamiento
Router8	3001::1/64			2008::1/64	RIPng

Vlans				
Switch	Modo de Configuración	N Vlans	Rango	Sist. Autónomo



Switch1	VTP Servidor	vlan10 VoIP: 192.168.50.0/24 Vlan12 administración: 192.168.100.0/24 Vlan 14 Telemática: 192.168.200.0/24 Vlan 16 Voz Datos: 192.168.5.0/24	fa 0/5-11 fa 0/12-16 fa 0/17-21 fa 0/22-24	1 1 1
Switch2	VTP Cliente	Definición del funcionamiento del VTP Cliente- Servidor		
Switch3	VTP Cliente	Definición del funcionamiento del VTP Cliente- Servidor		
Switch14		Vlan 6 Ingeniera: 3007:2:2::/64 Vlan 4 Docentes: 3007:1:1::/64	Fa 0/15-24 Fa 0/5-14	5 5
Switch6		Vlan8 Registro: 3006:1:1::/64 Vlan12 Datos: 3006:2:2::/64	Fa 0/5-14 Fa 0/15-24	5 5

DHCPv6 /Autoconfiguración				
Routers	Interfaz	Pool Name	Red	Sist. Autónomo



Router5 DHCPv6	0/1	Gerencia	2001:AB:2:2::/40 64	3
	1/0	Visitas	2020:DB:AB:88::/40 64	
Routers	Subinterfaces	Vlan	Dirección de Red	Sist. Autónomo
Router9 Auto	0/1.8	8 Registro	3006:1:1::1/64	5
	0/1.12	12 Datos	3006:2:2::1/64	
	0/0.4	4 Docentes	3007:1:1::1/64	
	0/0.6	6 Ingeniera	3007:2:2::1/64	

Equipos de Escritorios / Laptops /Servidor						
Switch2						
Nombre	Interfaz	Dinámico	Estático	Autoconfig	Rango/ Dirección IPv4	Vlan / AS
Laptop1	0/17	SI			192.168.200.0/24	Vlan: 14 (fa 0/17-21) AS: 1
Laptop2	0/18					
Laptop3	0/12	SI			192.168.100.0/24	Vlan: 12 (fa 0/12-16) AS 1
Laptop4	0/13					
Phone1	0/5	SI			192.168.50.0/24	



Phone2	0/6					Vlan: 10 (fa 0/5-11) AS 1
Phone3	0/23	SI			192.168.5.0/24	Vlan: 16 (fa 0/22-24) AS 1
Phone4	0/22					
PC1	PC	SI			192.168.5.0/24	Vlan: 16 (fa 0/22-24) AS 1
PC2	PC					
NOTA:						
- El Switch 2 configurado con VTP en modo Cliente. - El equipo PC1 está conectado directamente al Móvil Phone3. - El equipo PC2 está conectado directamente al Móvil Phone4. - El Switch 2 se conecta en modo Trunk con el Switch 1 a través de fa la 0/2.						
Switch3						
Nombre	Interfaz	Dinámico	Estático	Autoconfig	Rango/ Dirección IPv4	Vlan / AS
Laptop5	0/12	SI			192.168.100.0/24	Vlan: 12 (fa 0/12-16) AS 1
Laptop6	0/13					
Laptop7	0/17	SI			192.168.200.0/24	Vlan: 14 (fa 0/17-21) AS: 1
Laptop8	0/18					
Phone5	0/5	SI			192.168.50.0/24	Vlan: 10 (fa 0/5-11) AS 1
Phone6	0/6					



PC3	PC	SI			192.168.5.0/24	Vlan: 16 (fa 0/22-24) AS 1
PC4	PC					
Phone7	0/22	SI			192.168.5.0/24	Vlan: 16 (fa 0/22-24) AS 1
Phone8	0/23					

NOTA:

- El Switch 2 configurado con VTP en modo Cliente.
- El equipo PC4 está conectado directamente al Móvil Phone8.
- El equipo PC3 está conectado directamente al Móvil Phone7.
- El Switch 3 se conecta en modo Trunk con el Switch 1 a través de fa la 0/3

Switch1

- El Switch 1 se conecta con el router0 en modo Trunk a través de las interfaces 0/1.
- El Switch 1 fue configurado con VTP modo Servidor

Nombre	Interfaz	Dinámico	Estático	Autoconfig	Rango/ Dirección IPv4	Vlan / AS
ServDHCP-VLAN	0/1		Si		10.1.1.2/24	- AS 1

NOTA

En el servidor se crearon los siguientes pool: VoIP (Vlan 10), Administración (Vlan 12), Telemática (Vlan 14) y Voz Datos (Vlan 16).

Nombre	Interfaz	Dinámico	Estático	Autoconfig	Rango/ Dirección IPv4	Vlan / AS
--------	----------	----------	----------	------------	-----------------------	-----------



Laptop9	Fa 0/3	SI			2001:ab:2:2::/64	- AS 3
Laptop10	Fa 0/2					
Laptop11	Fa 0/2	SI			2020:db:ab:88::/64	- AS 3
Laptop12	Fa 0/3					
Laptop13	Fa 0/5			SI	3007:2:2::/64	Vlan 6 (fa 0/15-24) AS 5
Laptop14	Fa 0/15				3007:1:1::/64	Vlan 4 (fa 0/5-14) AS 5
Laptop15	Fa 0/5			SI	3006:1:1::/64	Vlan 8 (fa 0/5-14) AS 5
Laptop16	Fa 0/18				3006:2:2::/64	Vlan 12 (fa 0/15-24) AS 5
Laptop17	Fa 0/2		SI		3001::100/64	
Server1	Fa 0/1		SI		2020:1234:22:11::9/64	
Server4	Fa 1/1		SI		2008::2/64	

Aplicación del Nat64		
Dispositivos Finales	Dirección-original	Traducción
PC1	192.168.5.4/24	2200:: 192.168.5.4/96
PC2	192.168.5.3/24	2200:: 192.168.5.3/96
PC3	192.168.5.5/24	2200:: 192.168.5.5/96
PC4	192.168.5.2/24	2200:: 192.168.5.2/96



Laptop 0	2020:DB:AB:88:202:4AFF:FEE3:5684/64	23.23.23.3/24
Laptop 1	192.168.200.4/24	2200:: 192.168.200.4/96
Laptop 2	192.168.100.5/24	2200:: 192.168.100.5/96
Laptop 3	192.168.100.4/24	2200:: 192.168.100.4/96
Laptop 4	2001:AB:2:2:2E0:F9FF:FEBD:3E4B/64	24.24.24.2/24
Laptop 5	192.168.100.3/24	2200:: 192.168.100.3/96
Laptop 6	192.168.200.2/24	2200::192.168.200.2/96
Laptop 7	192.168.100.2/24	2200::192.168.100.2/96
Laptop 8	192.168.200.5/24	2200:: 192.168.200.5/96
Laptop 9	192.168.200.3/24	2200:: 192.168.200.3/24
Laptop 10	3007:1:1:0:206:2AFF:FEDE:E2C8/64	35.35.35.2/24
Laptop 11	3006:1:1:0:203:E4FF:FE18:2084/64	33.33.33.2/24
Laptop 12	3001::100/64	30.30.30.100/24
Laptop 13	2020:DB:AB:88:260:2FFF:FE2B:44D2/64	23.23.23.2/24
Laptop 14	2001:AB:2:2:201:C7FF:FEB8:7291/64	24.24.24.3/24
Laptop 17	3006:2:2:0:210:11FF:FE89:104C/64	32.32.32.2/24
Laptop 18	3007:2:2:0:260:5CFF:FE9A:4E0C/64	34.34.34.2/24
Phone1	192.168.50.6/24	2200:: 192.168.50.6/96
Phone2	192.168.50.11/24	2200:: 192.168.50.11/96
Phone3	192.168.50.7/24	2200:: 192.168.50.7/96
Phone4	192.168.50.10/24	2200:: 192.168.50.10/96
Phone5	192.168.50.2/24	2200:: 192.168.50.2/96
Phone6	192.168.50.3/24	2200:: 192.168.50.3/96
Phone7	192.168.50.4/24	2200:: 192.168.50.4/96
Phone8	192.168.50.5/24	2200:: 192.168.50.5/96
Server0	10.1.1.2/24	2200:: 10.1.1.2/96
Server1	2020:1234:22:11::9/64	26.26.26.9/24
Server4	2008::2/64	29.29.29.2/24

NOTA: Recuerde que si tiene problemas al momento de tratar de comunicarse con el dispositivo final, la recomendación sería de ingresar una dirección IP NAT64 para llegar al router al cual está conectado el equipo. Esta dirección NAT64 del Router debe coincidir con la dirección de red NAT64 del dispositivo final.



La dirección NAT64 tanto del router conectado al equipo, como del equipo debe de coincidir, es decir deben de estar los dispositivos dentro de la misma red NAT64.

Ejemplo: Laptop 1 su dirección IP es 192.168.200.4/24 su dirección NAT64 sería 2200::192.168.200.4/96
Router0 su dirección IP es 192.168.200.1/24 su dirección NAT64 sería 2200::192.168.200.1/96

ENUNCIADO

Asignación de direcciones IP.

Se deberá asignar las direcciones IP a las interfaces de los Routers tal y como aparecen en el cuadros llamados: **Direccionamiento IPv6 / Sistema Autónomo 1, Direccionamiento IPv6 / Sistema Autónomo 2, Direccionamiento IPv6 / Sistema Autónomo 3, Direccionamiento IPv6 / Sistema Autónomo 4, Direccionamiento IPv6 / Sistema Autónomo 5 y Direccionamiento IPv6 / Sistema Autónomo 6.**

Configuración de equipos estático.

Se deberá asignar las direcciones IP a las interfaces de Pc Escritorio y Laptops, tal como aparecen en el cuadro llamado: **Equipos de Escritorios / Laptops /Servidor.**

Configuración de VLANs.

Se deberá crear y configurar las VLANs, tal como aparecen en el cuadro llamado: **VLANs.**

Configuración DHCPv6 y Autoconfiguración.

Para la asignación de direcciones dinámicas se deberá crear los Pool tal y como se muestra en los siguientes cuadros: **DHCPv6 /Autoconfiguración.**

La autoconfiguración se realizaran en el siguientes Router: **Router9.** Los equipos finales que están conectado al Routers se le asignan direcciones IP por autoconfiguración.

Para realizar la Autoconfiguración, el único requisito es que la interface del Router, tenga asignado un prefijo de Red.



Configuración de protocolo de enrutamiento Dinámico Interno

Se deberá asignar el protocolo definido de los siguientes cuadros: **Direccionamiento IPv4 / Sistema Autónomo 1 y Direccionamiento IPv6 / Sistema Autónomo 1 hasta el 6**

Redistribución de rutas

La redistribución de rutas se realizara en los siguiente Routers:

- Routers => Router2: Redistribución de RIPng con EIGRPv6
- Routers => Router7: Redistribución de RIPng con EIGRPv6
- Routers => Router8: Redistribución de RIPng, EIGRPv6 y OSPFv3.

Configuración de Teléfono VOIP

Para la configuración de los teléfonos con VOIP, se creó una VLAN llamada VOIP (Con dirección de red: 192.168.50.0), por lo cual ya tiene asignadas las interfaces correspondiente, como se puede observar en los Switch 2 y 3

Configuración de NAT 64.

Para la configuración del Nat64, es requerido recordar lo que se establece es la traducción de dirección IPv4 a IPv6 y viceversa, aquí se usara NAT64 tradicional (de uno a uno) como se puede observar en el cuadro: **Aplicación NAT64.**

Tiempo estimado de solución

- 10 horas

Preguntas de análisis

En los siguientes apartados se pretende que los alumnos sean capaces de analizar y responder las siguientes preguntas con respecto al tema de **VOIP en IPv6**, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. ¿Por qué en VoIP se configura redes IPv4 y no IPv6?
2. ¿Es necesario que VoIP use algún mecanismo de transición de IPv4/IPv6? Justifique su respuesta



3. ¿Para qué sirve la siguiente línea de código **Option 150 IP 192.168.1.1** en VoIP?

4. Porque es necesario la creación de VLANs para el uso de VoIP?

5. ¿Qué indica las siguientes líneas de códigos?

```
telephony-service
max-ephones 10
max-dn 8
ip source-address 192.168.50.1 port 2000
auto assign 1 to 5
auto assign 5 to 8
```

6. En los enunciados elija la afirmación que crea más conveniente para el uso de VoIP:

- a) Es necesario el uso de NAT64 en VoIP ya que permite ver una red IPv4 como si fuera IPv6 y viceversa.
- b) No se necesita ningún tipo mecanismo de transición para comunicar VoIP en redes IPv6.
- c) El comando **ephone 1** añade un nuevo teléfono a la red VoIP con identificador 1,
- d) El inciso A y C es correcto.
- e) Ninguna de las anteriores.



PRÁCTICA 11: DNS

Objetivo general

- Implementar el funcionamiento de un servidor DNS en IPv6 usando CORE.

Objetivos específicos

- Establecer la configuración de un servidor DNS en el direccionamiento de red de IPv6.
- Analizar y comprender el funcionamiento de un servidor DNS usando capturas de tráfico en Wireshark.

Introducción

Es una tecnología basada en una base de datos que sirve para **resolver nombres** en las redes, es decir, para conocer la dirección IP de la máquina donde está alojado el dominio al que queremos acceder.

DNS es un sistema que sirve para traducir los nombres en la red, y está compuesto por tres partes con funciones bien diferenciadas como son el **Cliente DNS, Servidor DNS y Zonas de autoridad**.

Requerimientos

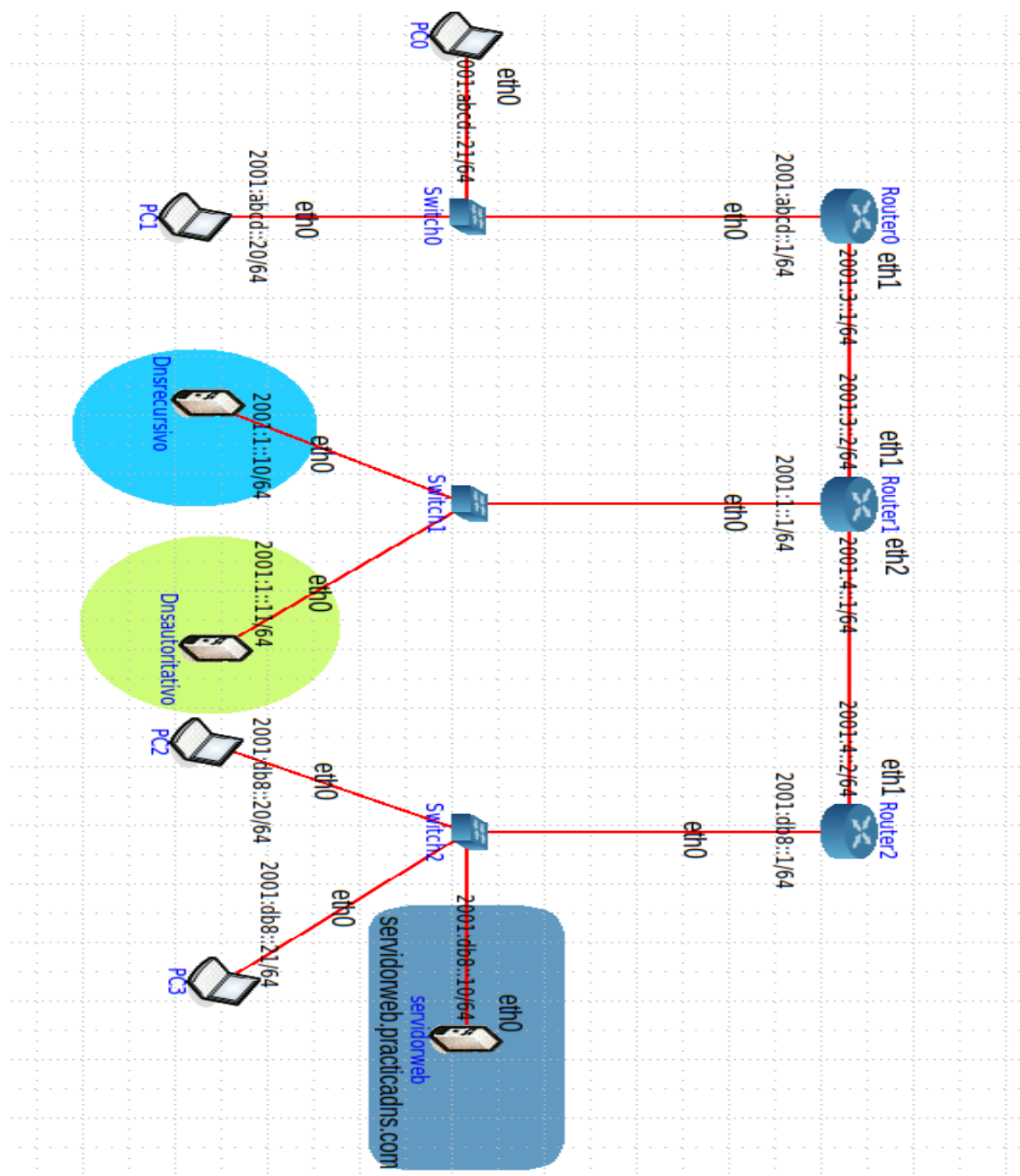
Hardware	Software
Computadora con los siguientes requisitos: ➤ Procesador mínimo de velocidad de 2.1 GHz ➤ Memoria RAM de 2 GB.	Emulador de redes CORE v4.7 con los siguientes elementos: ➤ 3 Router ➤ 3 Switches ➤ 4 PCs ➤ 3 servidores Nota: Se debe instalar Wireshark para la captura de tráfico.

Conocimientos previos

Para la correcta realización de esta práctica el estudiante deberá tener conocimientos básicos sobre configuración de DNS usando bind9, implementación de Quagga, uso de Wireshark para análisis de tráfico de redes. Además del uso de otras tecnologías y protocolos en el direccionamiento de IPv6.



Topología



Funcionalidad

En la figura de la práctica anterior se hace uso de servidor de nombres DNS donde las PCs que están conectadas a dicha red harán una solicitud a un dominio que está registrado en el mismo. Para la conexión de los diferentes dispositivos presentes en este escenario se hace uso del protocolo de enrutamiento dinámico RIPng. También para captura de tráfico se usara Wireshark para analizar el funcionamiento de un servidor DNS.



Comandos de ayuda

Comandos	Descripción
apt-get install autoremove purge paquete.deb	Instala, desinstala o remueve por completo un paquete.deb
service bind9 start stop restart /etc/init.d/bind9 start stop restart	Sirve para iniciar, detener o reiniciar el servicio bind9 (DNS).
ping6 direccion-ipv6 hostwithipv6address -I <device> <link-local-ipv6address>	Envía un ICMP ECHO_REQUEST to network hosts
named-checkconf	informa de posibles errores en el fichero /etc/bind/named.conf.local
named-checkzone <dominio> <fichero> Ejemplo : \$named-checkzone aula202.com /etc/bind/db.aula202	Comando que ayuda a encontrar errores en el fichero de registro de recursos.
nameserver dirección-ip	Traduce los nombres a direcciones dentro de una red.
dig [@servidor_dns] [opciones] [tipo]	Permite comprobar tanto el mapeo de nombres a IPs como el mapeo inverso de IPs a nombres, pero sólo sirve para Internet, ya que no mira en /etc/hosts (sólo utiliza /etc/resolv.conf).
host [-aCdIriTw] [-t type] [-R number] hostname [server]	El comando host se usa para encontrar la dirección IP del dominio dado y también muestra el nombre de dominio para la IP dada.

Datos de los dispositivos

Servidores			
Nombres	Interfaz	Dirección de red IPv6	Dirección IPv6
Dnsautoritativo	Eth0	2001:1::/64	2001:1::11/64
Dnsrecursivo	Eth0	2001:1::/64	2001:1::10/64
Servidorweb	Eth0	2001:db8::/64	2001:db8::10/64
Router0			
Eth0		2001:abcd::/64	2001:abcd::1/64
Eth1		2001:3::/64	2001:3::1/64
Router1			
Eth0		2001:1::/64	2001:1::1/64



Eth1	2001:3::/64	2001:3::2/64
Eth2	2001:4::/64	2001:4::1/64
Router3		
Eth0	2001:db8::/64	2001:db8::1/64
Eth1	2001:4::/64	2001:4::2/64

Clientes			
Dispositivo	Interfaz	Dirección de red	Dirección ipv6
PC0	Eth0	2001:abcd::/64	2001:abcd::21/64
PC1	Eth0	2001:abcd::/64	2001:abcd::20/64
PC2	Eth0	2001:db8::/64	2001:db8::20/64
PC3	Eth0	2001:db8::/64	2001:db8::21/64

ENUNCIADO

Asignación de direcciones IP

Asigne la dirección a cada uno de los dispositivos tal y como se muestra en las tablas anteriores donde se detalla cada dirección que usa cada equipo y la interfaz que tiene conectada. Para saber si se ha configurado correctamente el escenario puede hacer uso del comando **ping6 dirección-ip**.

NOTA: Si tiene problemas de respuesta por parte de una dirección determinada usando el comando **ping6 dirección-ip**, revise el archivo de configuración en el directorio `/etc/quagga/Quagga` con el visor o editor de texto que prefiera.

Protocolo de encaminamiento dinámico

El protocolo de enrutamiento dinámico que se usó es RIPng, este se usó en todo el escenario de la red mostrada en la imagen. Recuerde que puede habilitar este protocolo usando un Gui (interfaz gráfica) que nos ofrece CORE y/o también usando un archivo de configuración como es el siguiente: `/etc/quagga/daemons`.



Servidores DNS

Para el funcionamiento del servidor DNS se debe instalar bind9 desde el terminal o de forma manual. Los servidores DNS estarán compuesto por los siguientes:

➤ DNS Autoritativo:

En el equipo que se llama dnsautoritativo se deberá de crear una zona con el nombre **practicadns.com** que responderá tanto al nombre como a su dirección inversa. La zona se creara en el archivo de configuración **/etc/bind/named.conf.default.local** y los ficheros que también se configurarán son los siguientes **/etc/bind/db.practicadns** y **/etc/bind/db.practicaipv6reverse** donde se registrara el nombre de dominio respectivamente.

➤ DNS recursivo

Recuerde que el servidor DNS recursivo solo reenvía la consulta hacia otros servidores que se le indique en el archivo de configuración **/etc/bind/named.conf.option**. Este modo de uso se implementara en el dispositivo que se llama dnsrecursivo.

➤ Creación de un subdominio

El subdominio que va a crearse se llama servidorweb.practicadns.com que corresponde con el nombre del equipo llamado servidorweb. El subdominio se debe de ingresar en los archivos de configuración **/etc/bind/db.practicadns** y **/etc/bind/db.practicaipv6reverse**.

NOTA: Recuerde que debe de agregar la dirección IP del servidor DNS en el archivo de configuración **/etc/resolv.conf** usando **nameserver ip-dns** para la resolución de nombres de los dominio registrados en la red.

Análisis de tráfico

Para la verificación de las consultas y respuestas por parte de un servidor DNS frente a maquinas clientes se hará uso de Wireshark para analizar el funcionamiento del servicio DNS. Las capturas de tráfico con Wireshark se implementan en la PC0 u otro cualquier Pc que crea conveniente.

Tiempo estimado de solución

- 7 horas.



Preguntas de análisis.

En los siguientes apartados se pretende que los estudiantes sean capaces de analizar y responder las siguientes preguntas con respecto al tema de DNS, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. ¿Qué es un servidor DNS autoritativo?
2. ¿Es necesario la creación de un servidor DNS recursivo? justifique su respuesta Sí, No.
3. ¿Cuál es la diferencia entre un servidor DNS autoritativo y uno recursivo?
4. ¿Por qué es importante agregar el servidor de nombre nameserver en el archivo
`/etc/resolv.conf`?
5. ¿Explique la ventaja del comando `host` vs `nslookup`?
6. ¿Cuál es la importancia de analizar el tráfico de paquetes de un servidor DNS en la red?
7. En la captura siguiente que se realizó en la PC0 ¿Explique los pasos que se llevan a cabo desde que la PC0 hace la solicitud hasta el momento que responde el Servidor DNS?

No.	Time	Source	Destination	Protocol	Length	Info
1	0.000000000	2001:abcd::21	2001:1::11	DNS	107	Standard query 0xd334 A servidorweb.practicadns.com
2	0.000467000	2001:1::11	2001:abcd::21	DNS	164	Standard query response 0xd334
3	2.001210000	2001:abcd::21	2001:1::11	DNS	107	Standard query 0xe168 AAAA servidorweb.practicadns.com
4	2.001521000	2001:1::11	2001:abcd::21	DNS	177	Standard query response 0xe168 AAAA 2001:db8::10
5	2.525573000	fe80::200:ff:feaa:0	ff02::9	RIPng	166	Command Response, Version 1
6	4.002267000	2001:abcd::21	2001:1::11	DNS	107	Standard query 0x9f04 MX servidorweb.practicadns.com
7	4.002546000	2001:1::11	2001:abcd::21	DNS	164	Standard query response 0x9f04

Frame 1: 107 bytes on wire (856 bits), 107 bytes captured (856 bits) on interface 0

Ethernet II, Src: 00:00:00_aa:00:0d (00:00:00:aa:00:0d), Dst: 00:00:00_aa:00:00 (00:00:00:aa:00:00)

Internet Protocol Version 6, Src: 2001:abcd::21 (2001:abcd::21), Dst: 2001:1::11 (2001:1::11)

User Datagram Protocol, Src Port: 60334 (60334), Dst Port: 53 (53)

Source Port: 60334 (60334)

Destination Port: 53 (53)

Length: 53

Checksum: 0xa07e [validation disabled]

[Good Checksum: False]

[Bad Checksum: False]

[Stream index: 0]

Domain Name System (query)



PRÁCTICA 12: HTTP

Objetivo general

- Presentar el funcionamiento del servidor HTTP con direcciones IPv6.

Objetivos específicos

- Verificar la correcta funcionalidad del servidor HTTP.
- Realizar pruebas de peticiones desde los clientes al servidor.

Introducción

En la siguiente práctica estudiaremos el funcionamiento de dicho servidor, el servidor que se utilizara es Apache2 para la demostración que es uno de los servicios más usado en la web y tiene soporte para Ipv6. Este servidor no se encuentra instalado por lo que será necesario descargarlo e instalarlo en el emulador.

Requerimientos:

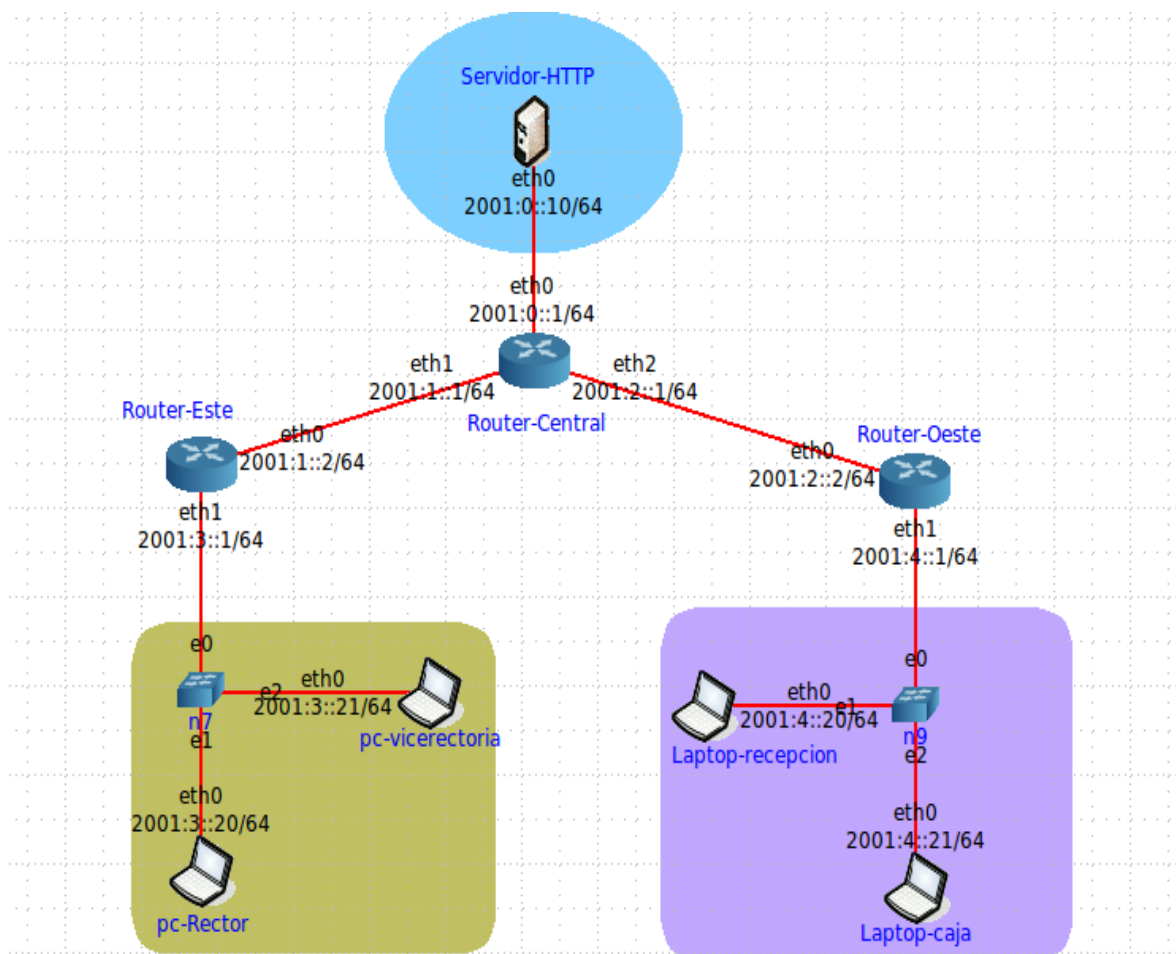
Hardware	Software
Computadora con los siguientes requisitos: ➤ Procesador con velocidad de 2.1 ghz ➤ Memoria RAM de 2 GB.	Emulador de redes Core con los siguientes elementos: ➤ 3 Router ➤ 3 Switches ➤ 4 PCs ➤ 1 servidor Nota: Instalar Wireshark para la captura de tráfico.

Conocimientos previos

Para la correcta realización de esta práctica el estudiante deberá tener conocimientos básicos de la configuración de Apache.



Topología



Funcionalidad

En la figura de esta práctica se muestra la topología en la cual se estará realizando la verificación y pruebas de dicho servidor con las maquinas clientes mostradas en la imagen.

Para iniciar con la práctica se recomienda seguir paso a paso lo siguiente:

- Instalar Wireshark.
- Se deberá instalar el servidor apache2.
- Iniciar dicho servidor luego de la instalacion.
- Verificar que el servidor este activado.
- Realizar prueba desde los clientes al servidor.
- Tomar captura de tráfico con Wireshark.

**Comandos de ayuda**

Comando	Definición
sudo apt-get install nom-paquete	Descarga e instala el paquete especificado.
/etc/init.d/apache2 [start stop]	Comando para iniciar o detener el servidor apache.
ps aux	Verifica que el proceso apache2 este activo.
netstat -antup	Verifica el puerto por el que está escuchando el servidor.
wget http:// [direccion-servidor]/	Comando para crear una solicitud del cliente al servidor http.
tcpdump -i eth0 -s 0 -w /tmp/nombre-captura_e1.pcap	Comando para realizar la captura de paquetes en Wireshark desde la consola.

Datos de los dispositivos

Servidor-HTTP		
Interfaz	Dirección de red IPv6	Dirección IPv6
Eth0	2001:0::/64	2001:0::10/64
Router-central		
Eth0	2001:0::/64	2001:0::1/64
Eth1	2001:1::/64	2001:1::1/64
Eth2	2001:2::/64	2001:2::1/64
Router-este		
Eth0	2001:1::/64	2001:1::2/64
Eth1	2001:3::/64	2001:3::1/64
Router-oeste		
Eth0	2001:2::/64	2001:2::2/64
Eth1	2001:4::/64	2001:4::1/64
Datos de las Pc.		



Dispositivo	Interfaz	Dirección de red	Dirección ipv6
Pc-Rector	Eth0	2001:3::/64	2001:3::20/64
Pc-Vicerrectoría	Eth0	2001:3::/64	2001:3::21/64
Laptop-recepción	Eth0	2001:4::/64	2001:4::20/64
Laptop-caja	Eth0	2001:4::/64	2001:4::21/64

ENUNCIADO

Asignación de direcciones:

Asignar las direcciones como se muestra en las tablas mencionadas en la práctica, para la secuencia de esta misma.

Protocolo de encaminamiento dinámico

El protocolo de enrutamiento dinámico que se uso es RIPng, este se usó en todo el escenario de la red mostrada en la imagen.

Para el funcionamiento del servidor, se debe instalar apache2.

- Para realizar las verificaciones entre el servidor y los clientes, será necesaria la utilización de del programa Wireshark, que mejora la visualización de paquetes transmitidos en la red.
- Una vez instalados los paquetes requeridos realizar lo siguiente:
- Iniciar el servicio apache2 desde el servidor.
- Desde la maquina cliente abrir una consola para capturar los paquetes, luego abrir otra consola del mismo cliente para iniciar la solicitud al servidor mientras se realiza la captura.

NOTA: No cierre la terminal donde se está ejecutando la captura hasta después de agregar el comando de petición del cliente al servidor, ya que puede ocasionar la interrupción del comando "tcpdump" y perjudicaría el resultado.

Después de haber realizado los incisos anteriores debe visualizar desde el cliente el archivo **index.html** para ver si se guardó correctamente.

Visualizar desde el servidor los logs generados por apache (/var/log/apache2/Access.log). También visualizar desde la captura de todo el proceso Wireshark.



Al final detener el servicio apache.

Tiempo de duración de la práctica

- 4 horas.

Preguntas de análisis

1. ¿Cómo se realiza una solicitud al servidor desde un cliente con direcciones IPv6?
 - a. # wget http://[2001.db8.10]/
 - b. # wget http://[2001:db8::10]/
 - c. # wget http://2001:db8::10/
 - d. # get http://[2001:db8::10]/

Las siguientes preguntas serán respondidas conforme a la siguiente imagen que muestra la captura de paquetes http.

The image shows a Wireshark capture of an HTTP GET request. The packet list shows two packets: packet 16 (198 bytes) and packet 18 (635 bytes). Packet 16 is the GET request, and packet 18 is the response. The packet details for packet 16 show the following structure:

- Frame 16: 198 bytes on wire (1584 bits), 198 bytes captured (1584 bits)
- Ethernet II, Src: 00:00:00:aa:00:09 (00:00:00:aa:00:09), Dst: 00:00:00:aa:00:06 (00:00:00:aa:00:06)
- Internet Protocol Version 6, Src: 2001:3::21 (2001:3::21), Dst: 2001::10 (2001::10)
- Transmission Control Protocol, Src Port: 40851 (40851), Dst Port: 80 (80), Seq: 1, Ack: 1, Len: 112
- Hypertext Transfer Protocol
 - GET / HTTP/1.1\r\n
 - User-Agent: wget/1.13.4 (linux-gnu)\r\n
 - Accept: */*\r\n
 - Host: [2001:0::10]\r\n
 - Connection: Keep-Alive\r\n
 - \r\n
 - [Full] request URI: http://[2001:0::10]/
 - [HTTP request 1/1]
 - [Response in frame: 18]

The packet bytes section shows the raw data of the request, including the GET method, HTTP version, and headers.



2. ¿Qué indica las dos tramas de HTTP?
 - a. La dirección 2001:3::21 está solicitando mediante un GET a la dirección 2001::10.
 - b. La dirección 2001::10 responde mediante un GET a la dirección 2001:3::21
 - c. La dirección 2001::10 responde mediante un ok a la dirección 2001:3::21.
 - d. Los incisos a y c son correctos.
3. ¿Qué protocolo se utiliza en IPv6 para la transmisión de datos entre el servidor HTTP y el cliente?
4. ¿En qué puerto escucha el servidor HTTP a los clientes con direcciones IPv6?



PRÁCTICA 13: INTÉRPRETE DE ÓRDENES SEGURA (SSH).

Objetivo general:

- Implementar SSH en redes IPv6 a través de Core Emú.

Objetivos específicos:

- Establecer conexiones seguras a través de SSH dentro de la red IPv6.
- Analizar mediante una captura de Wireshark como se logra establecer la conexión entre dos sistemas usando la arquitectura cliente/Servidor con SSH.

Introducción

La intérprete de órdenes segura conocida como SSH, tiene como finalidad acceder a un equipo a través de la red, y cumple con las siguientes funciones:

- Administrar Servidor remoto.
- Transferencia de archivos desde un ordenador remoto.
- Conexiones a través de LAN o Internet.
- Conexiones seguras y rápidas.
- Administración total de ordenador.

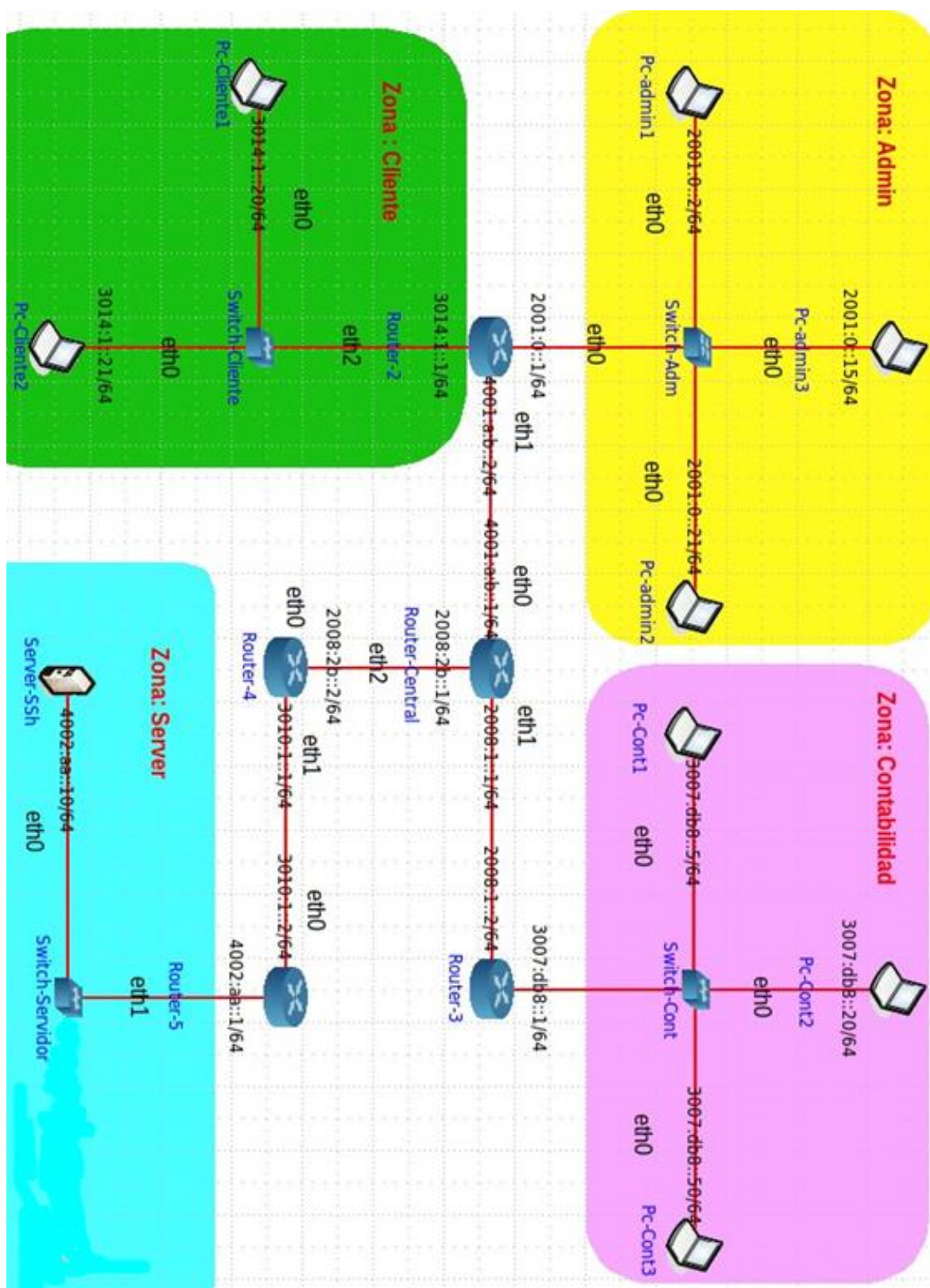
Requerimientos

Hardware	Software
➤ Computadora con los siguientes requisitos: • Procesador mínimo de velocidad de 2.1 GHz • Memoria RAM de 1 GB.	➤ Emulador de redes CORE v4.7 con los siguientes elementos: • 5 Router. • 4 Switchs. • 8 Pc Escritorio • 1 Servidor SSH.

Conocimientos previos: Para la correcta realización de esta práctica el alumno debe tener conocimientos teóricos y práctico de SSH en redes IPv4, aplicado en Linux-Ubuntu.



Topología:



**Funcionalidad:**

En la figura anterior se creó una topología con la finalidad de implementar el protocolo de transferencia de archivo y el intérprete de órdenes seguro dentro de un esquema de IPv6.

La topología está estructurada de la siguiente manera:

- **Zona Admin:** Pertenece a los equipos administrativos.
- **Zona Cliente:** Pertenece a los equipos que se le ofrece a los cliente para atender sus peticiones.
- **Zona Contabilidad:** Pertenece a los equipos del área de contabilidad.
- **Zona Server:** Es el área donde están instalado los servidores como SSH.

Además:

- Todo el esquema de red está configurado con IPv6.
- Se crearon los siguientes usuarios para implementarlo con SSH:
 - Red1: Hugo García.
 - Red2: William Zapata.
 - Red3: Delia Toruño.
- Se instaló los siguientes paquete/programas:
 - Apt-get install openssh-server.
 - Apt-get install Wireshark.
- El protocolo de enrutamiento aplicado fue RIPng en todo el esquema de red IPv6.

NOTA: La implementación de la práctica se realizó en Ubuntu 14.10 y para la instalación del Core Emú fue requerido los siguientes paquetes:

- sudo apt-get install core-network.
- sudo apt-get install quagga.
- sudo apt-get install bash bridge-utils ebtables iproute libtk-img python tcl8.5 tk8.5 autoconf automake gcc libev-dev make pkg-config python-dev libreadline-dev imagemagick help2man.

**Comandos de ayuda**

Comando o Acción	Descripción
ssh usuario@direccionIPv6 Example: ssh red1@3007:db8::50	Acceder a un sistema remoto a través de SSH
scp nombre_archivo usuario@[direccionIPv6]:/home/usuario/directorio/nombre_archivo Example: scp nombre_archivo red1@[3007:db8::50]: /home/red1/Documentos/nombre_archivo	Trasladar un archivo de modo seguro a través de SSH de un sistema remoto a otro.

Direccionamiento IPv6			
RIPNG			
Routers	Eth0	Eth1	Eth2
Router-Central	4001:a:b::1/64	200:1::1/64	2008:2b::1/64
Router-2	2001:0::1/64		3014:1::1/64
Router-3	2000:1::2/64	3007:db8::1/64	
Router-4	3010:1::1/64		2008:2b::2/64
Router-5	3010:1::2/64	4002:aa::1/64	

Equipos de Escritorios / Laptops /Servidor		
Nombre	Interfaz	Dirección IP
Pc-Admin1	ETH0	2001:0::2/64
Pc-Admin2	ETH0	2001:0::21/64
Pc-Admin3	ETH0	2001:0::15/64
Pc-Cont1	ETH0	3007:db8::5/64
Pc-Cont2	ETH0	3007:db8::20/64
Pc-Cont3	ETH0	3007:db8::50/64



Pc-Cliente1	ETH0	3014:1::20/64
Pc-Cliente2	ETH0	3014:1::21/64
Server-SSH	ETH0	4002:aa::10/64

ENUNCIADO

Asignación de direcciones IP

Asignar la dirección a cada uno de los dispositivos tal y como se muestra en las tablas anteriores donde se detalla cada dirección que usa cada equipo y la interfaz que tiene conectada. Para saber si se ha configurado correctamente el escenario puede hacer uso del comando **ping6 dirección-IPv6**.

Protocolo de Enrutamiento dinámico.

El protocolo de enrutamiento dinámico que se uso es RIPng en todo el escenario, recuerde que puede habilitar este protocolo usando un Gui (interfaz gráfica) que nos ofrece CORE y/o también usando un archivo de configuración como es: **/etc/quagga/daemons**.

Servidor SSH

Se configuro un servidor SSH, con la finalidad de conectarse con otro usuario de forma segura mediante la red y a la vez ser capaz de transferir todo tipo de archivos entre los usuarios. Para la correcta aplicación de los comando de conexión con usuario y transferencia de archivos visualizar el cuadro llamado: **Comando de ayuda**.

Tiempo estimado de solución:

- 4 horas.



Mediante la siguiente captura de Wireshark responda las siguientes preguntas:

Filter:		Expression...		Clear	Apply	Guardar
No.	Time	Source	Destination	Protocol	Length	Info
1	0.000000000	fe80::200:ff:feaa:0	ff02::9	RIPng	86	Command Response, Version 1
2	2.647624000	2001::2	3007:db8::20	TCP	94	50568->22 [SYN] Seq=0 Win=28800 Len=0 MSS=1440 SACK_PERM=1 TSval=188591 TSecr=0
3	2.647926000	3007:db8::20	2001::2	TCP	94	22->50568 [SYN, ACK] Seq=0 Ack=1 Win=28560 Len=0 MSS=1440 SACK_PERM=1 TSval=188591 TSecr=188591
4	2.647962000	2001::2	3007:db8::20	TCP	86	50568->22 [ACK] Seq=1 Ack=1 Win=28800 Len=0 TSval=188591 TSecr=188591
5	2.648216000	2001::2	3007:db8::20	SSHv2	120	Client: Protocol (SSH-2.0-openssh 6.6.1p1 ubuntu-8)
6	2.648257000	3007:db8::20	2001::2	TCP	86	22->50568 [ACK] Seq=1 Ack=35 Win=28672 Len=0 TSval=188591 TSecr=188591
7	2.654313000	3007:db8::20	2001::2	SSHv2	120	Server: Protocol (SSH-2.0-openssh 6.6.1p1 ubuntu-8)
8	2.654450000	2001::2	3007:db8::20	TCP	86	50568->22 [ACK] Seq=35 Ack=35 Win=28800 Len=0 TSval=188592 TSecr=188592
9	2.655261000	2001::2	3007:db8::20	SSHv2	2054	Client: Key Exchange Init
Transmission Control Protocol, Src Port: 50568 (50568), Dst Port: 22 (22), Seq: 1, Ack: 1, Len: 34						
Source Port: 50568 (50568)						
Destination Port: 22 (22)						
[Stream index: 0]						
[TCP Segment Len: 34]						
Sequence number: 1 (relative sequence number)						
[Next sequence number: 35 (relative sequence number)]						
Acknowledgment number: 1 (relative ack number)						
Header Length: 32 bytes						
..... 0000 0001 1000 = Flags: 0x018 (PSH, ACK)						
0000	00 00 00 aa 03 00 00	00 aa 00 00 86 dd 60 00				
0010	00 00 00 42 06 40 20 01	00 00 00 00 00 00 00 00	...B.@.....			
0020	00 00 00 00 00 02 30 07	0d b8 00 00 00 00 00 00	0.....			
0030	00 00 00 00 00 20 c5 88	00 16 84 96 b8 f5 dd b9	*			
0040	03 10 89 18 00 e1 5e 2a	00 00 01 01 08 0a 00 02	SS H-2-0-p			
0050	e0 af 00 02 e0 af 53 53	48 2d 32 2e 30 2d 4f 70	SSH 6. 6.1p1 ub			
0060	65 6e 53 53 48 5f 36 2e	36 2e 31 70 31 20 55 62	untu-8..			
0070	75 6e 74 75 2d 38 0d 0a					



Preguntas:

1. Explique ¿Qué indica la quinta línea de código en la imagen anterior?
2. ¿Porque es necesario el uso de conexión segura TCP?
3. ¿Es SSH un modelo de conexión cliente, servidor? Justifique
4. ¿Cuál es el procedimiento que se debe de seguir para iniciar sesión vía SSH?
5. Identifique en qué momento se realiza la conexión SSH y explique los campos que se visualizan en la captura.



PRÁCTICA 14: MISCELANEA IPv6

Objetivo general:

- Demostrar las múltiples tecnologías desarrolladas para llevar a cabo la ejecución del protocolo IPv6 en el entorno de redes.

Objetivos específicos:

- Determinar y establecer el proceso de configuración de cada uno de los protocolos de enrutamiento dentro de los diferentes sistemas autónomos creados (IPv6 e IPv4).
- Definir e implementar el método de distribución de rutas en los dispositivos de enrutadores con el fin de comunicar y asociar los diferentes protocolos de enrutamiento (IPv6 e IPv4).
- Implementar los mecanismos de transición, para lograr la comunicación entre IPv4 e IPv6.
- Configurar VoIP con la finalidad de establecer una llamada mediante IPv6.
- Implementar la lista de acceso con la finalidad de permitir o denegar un servicio a un determinado equipo o red.
- Aplicar los servicios de HTTP (**Hypertext Transfer Protocol**), FTP (**File Transfer Protocol**), DNS (Domain Name System) en IPv6 e IPv4.
- Implementar DHCP tanto para IPv6 como IPv4.
- Implantación de VLANs en IPv6 e IPv4.

Introducción:

IPv6 es la siguiente generación de estándares de protocolos de internet para direcciones IP, su objetivo es la de garantizar la seguridad, estabilidad y crecimiento del internet, por lo cual es necesario tener conocimientos concreto de esta generación de protocolo.

El desarrollo de esta práctica tiene como finalidad exponer innumerables herramientas que se pueden desarrollar bajo este protocolo, entre las cuales destacamos: Mecanismo de transición, direccionamiento dinámico y estático de IPv6, servicio HTTP, FTP, DNS entre otros.

Todo lo antes mencionado se verá reflejado tanto en la práctica como en el documento, en el cual se explica los pasos a seguir para la correcta configuración y entendimientos de cada herramienta aplicada a esta generación de protocolo.

**Requerimientos:**

Hardware	Software
Computadora con los siguientes requisitos: • Procesador mínimo de velocidad de 2.1 GHz • Memoria RAM de 1 GB.	Simulador Packet Tracer 6.1.1 : • 44 Laptops. • 25 Router 2811. • 22 Switch 2960. • 6 Server. • 5 VOIP. • 2 PC Escritorio. • Nube Frame Relay

Conocimientos previos:

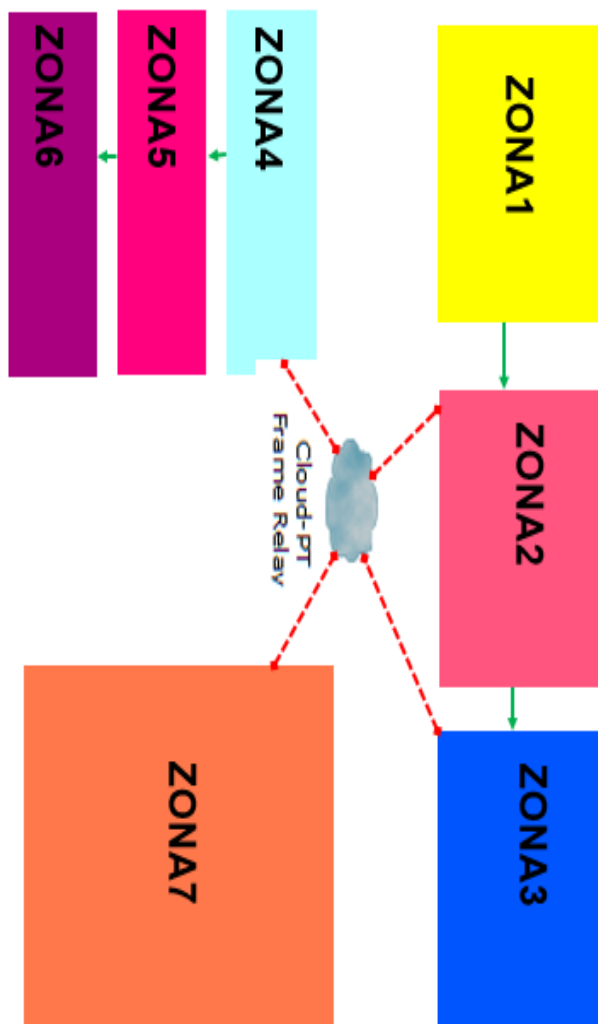
Para la correcta realización de esta práctica el estudiante deberá tener conocimientos:

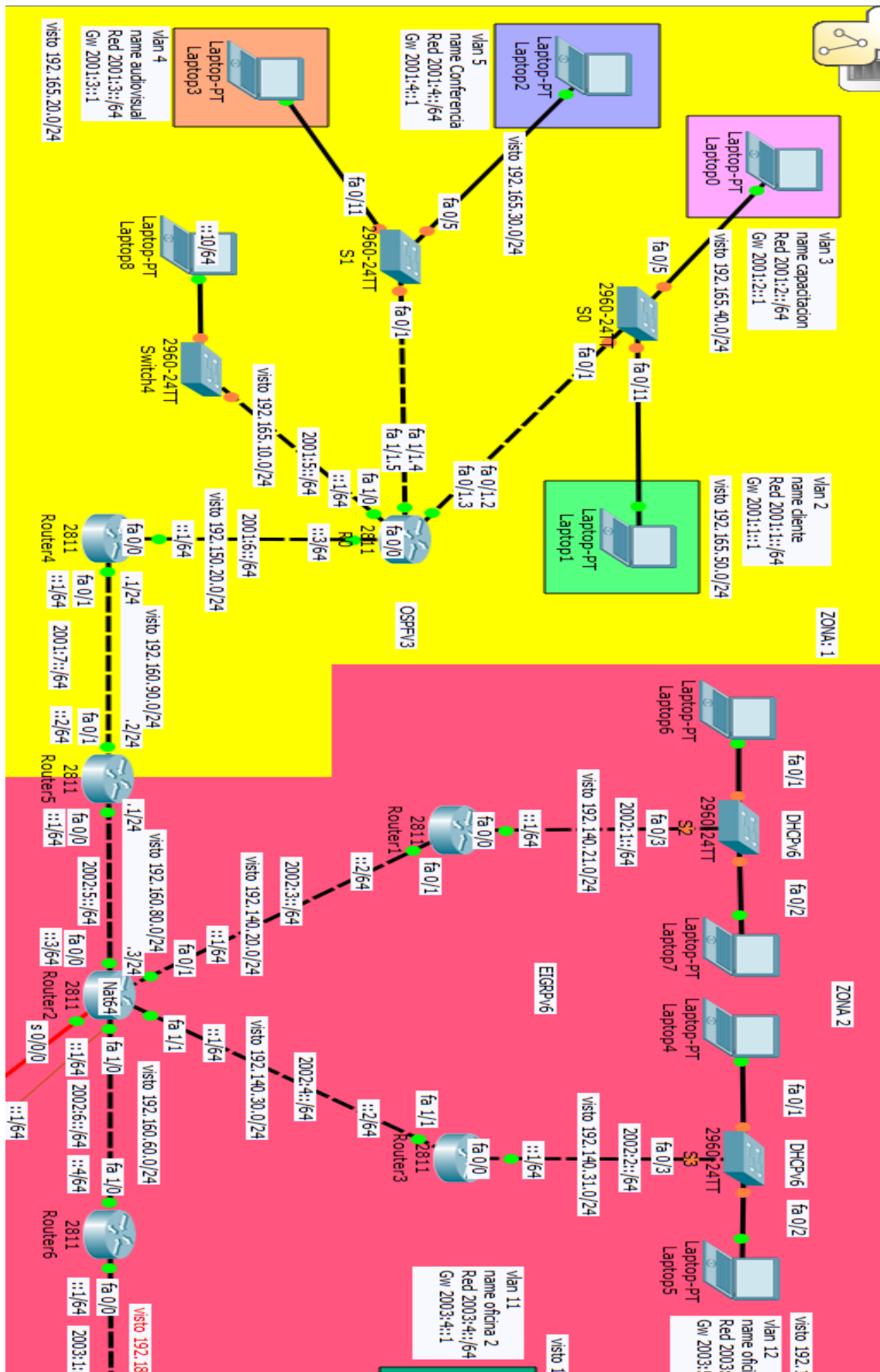
IPv4	IPv6
• Subnetting IPv4. • DHCP • Direccionamiento estático y dinámico (Ripv2, OSPF, EIGRP). • Redistribución de rutas entre los diferentes protocolos de enrutamiento dinámico. • Configuración de VLANs y VTP. • VOIP. • Configuración de Servidores DNS, HTTP y FTP. • Frame Relay • Listas de Acceso • NAT-PAT	• Subnetting IPv6. • DHCPv6 y Autoconfiguración. • Direccionamiento estático y dinámico (RIPng, OSPFv3, EIGRPv6). • Redistribución de rutas entre los diferentes protocolos de enrutamiento dinámico. • Configuración de VLANs. • Mecanismos de Transición. • Configuración de Servidores DNS, HTTP y FTP. • Listas de Acceso • NAT64

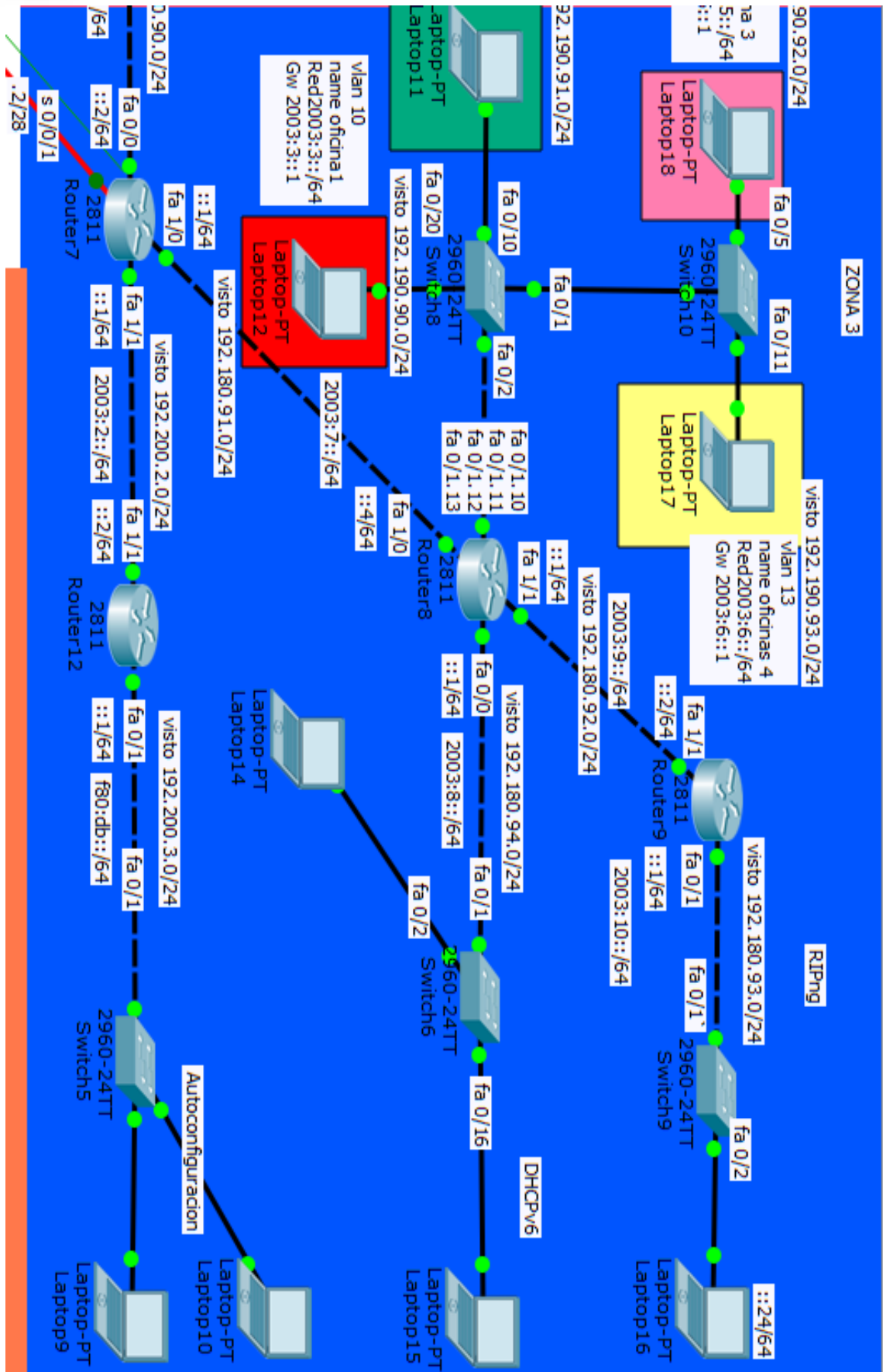


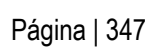
Topología:

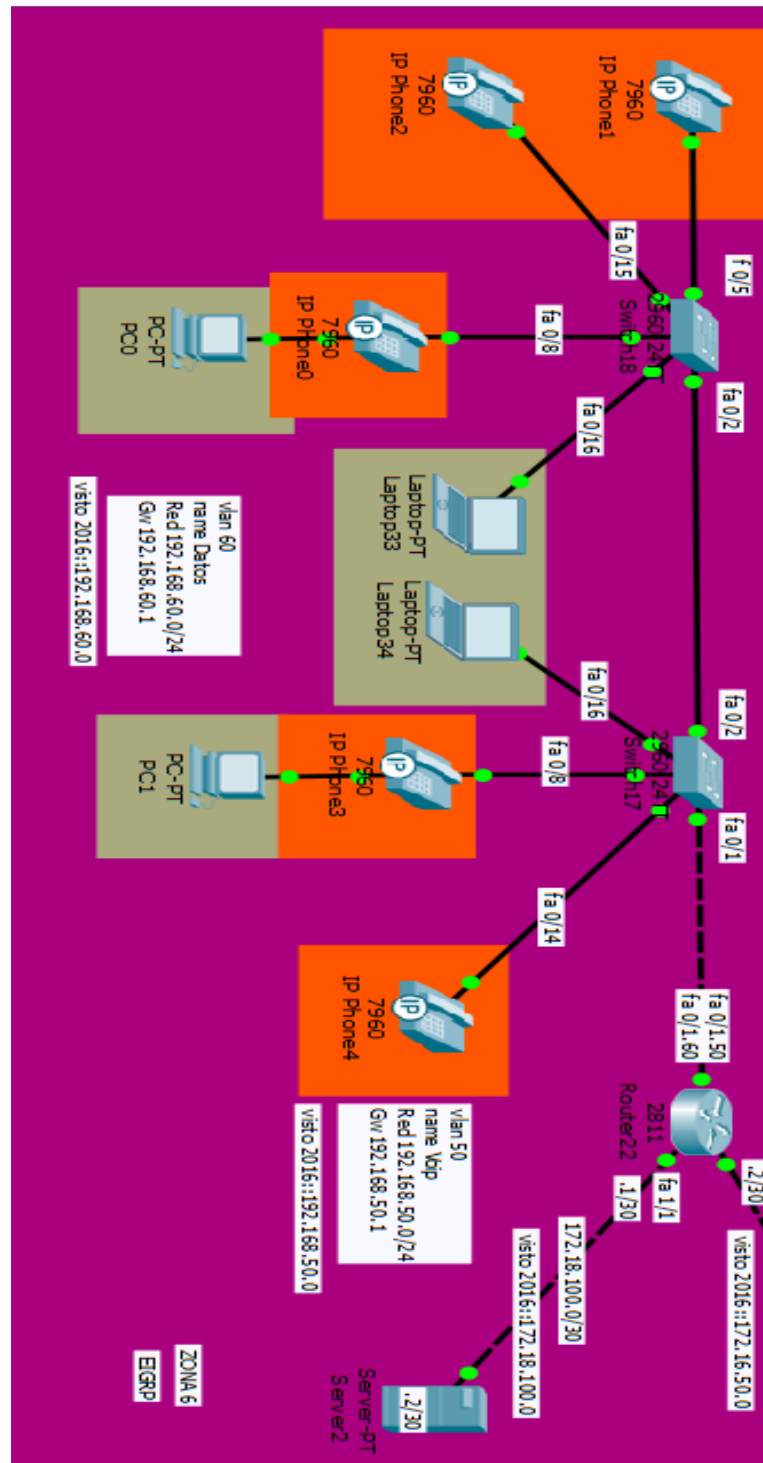
Esquema de la Miscelánea

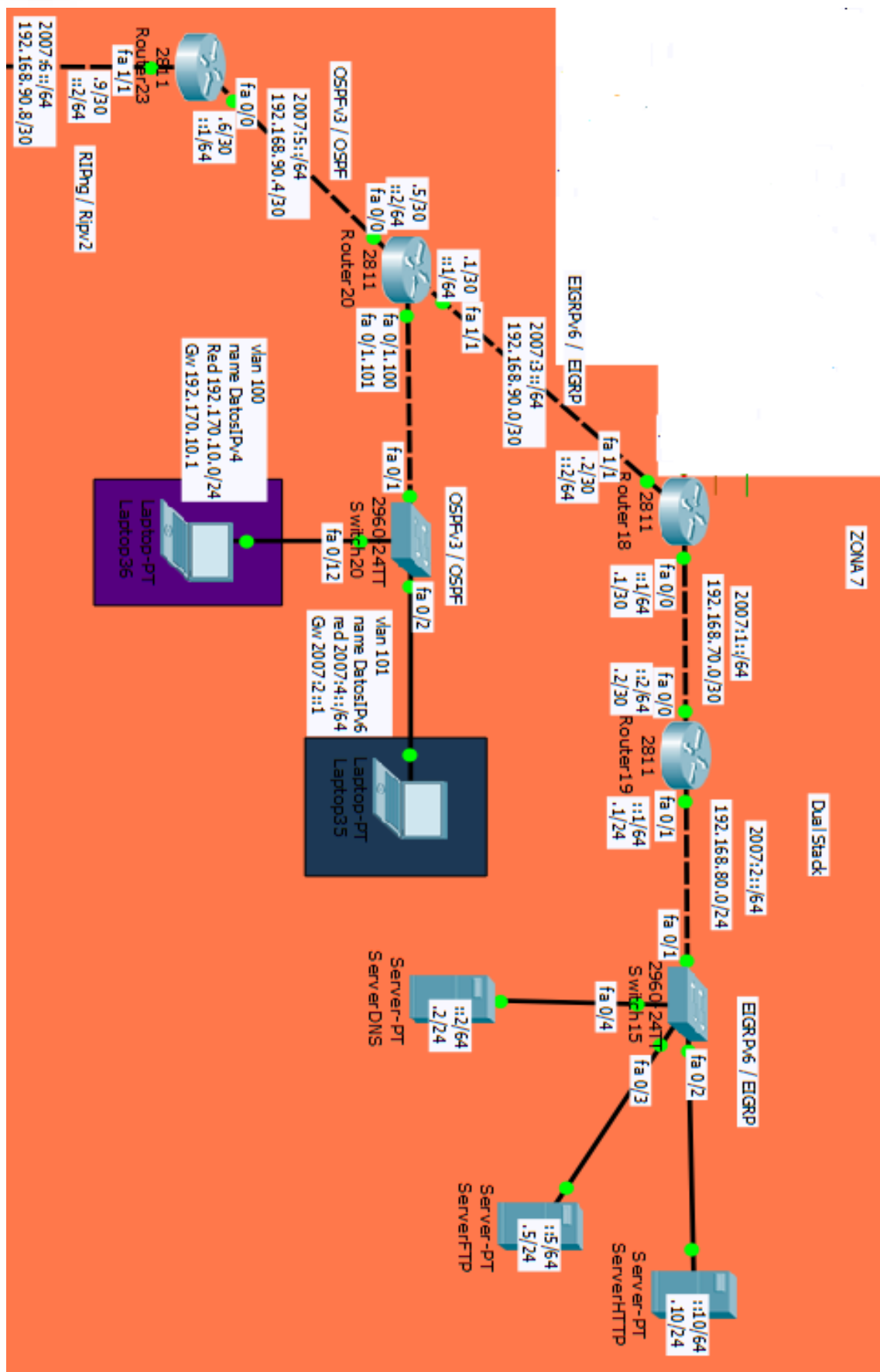


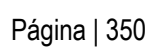






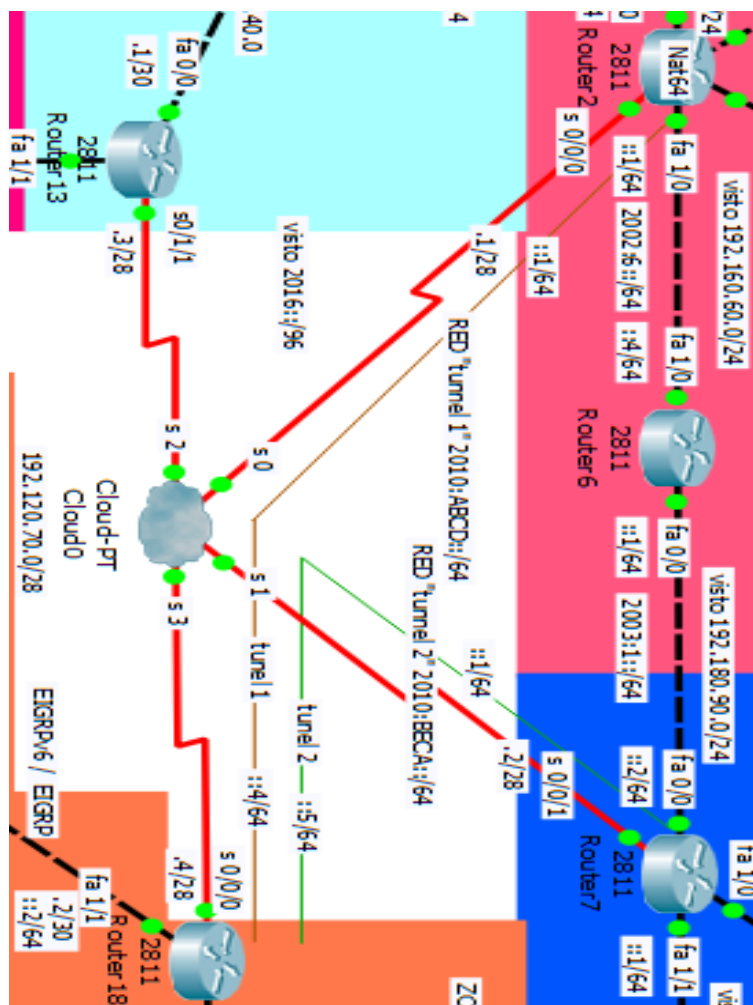








Frame Relay



Observación:

Aquí se detalla los routers que están conectados con la nube Frame Relay.

- Zona 1 -> Router2.
- Zona 2 -> Router7.
- Zona 4 -> Router13.
- Zona 7 -> Router18.

**Funcionalidad:**

En la figura anterior se identifica un esquema de red, en el cual se implementó IPv6 e IPv4, con la finalidad de exponer el funcionamiento e interacción entre ambos protocolos.

La topología está estructurado de la siguiente manera:

Zona 1: Se configuro el protocolo de enrutamiento dinámico OSPFv3, el cual está confeccionado con los siguientes routers: R0, R4 y R5. El área a que pertenecerán las redes será el Área "0". Se crearon las VLANs denominadas: Capacitación, Cliente, Conferencia y Audiovisual, la cuales recibirán direcciones IP mediante el uso DHCPv6 y esta estarán conectada al Router0.

Zona 2: Se configuro el protocolo de enrutamiento dinámico EIGRPv6, el cual está confeccionado con los siguientes routers: Router1, Router2, Router3, Router5 y Router6. Los equipos finales conectados a la red recibirán direcciones IP mediante el uso del DHCPv6.

Zona 3: Se aplicó el protocolo de enrutamiento dinámico RIPng, el cual está confeccionado con los siguientes routers: Router6, Router7, Router8, Router9 y Router12. Se crearon las siguientes VLANs denominadas: Oficina1, Oficina2, Oficinas3 y Oficinas4, las están conectada al Routers8. Los equipos conectados a esta VLANs recibirán sus respectivas direcciones IP de forma dinámica mediante el DHCPv6, la configuración de las VLANs se realizó mediante VTP. En el Router12 los equipos finales se conectaran a la red en forma de Autoconfiguración y en el router9 los equipo finales se le configuro de forma estática.

Zona 4: Se aplicó el protocolo de enrutamiento dinámico OSPFv2, el cual está confeccionado con los siguientes routers: Router13, Router14 y Router15. El área a que pertenecerán las redes será Área "0". Los equipos finales conectados a la red recibirán las direcciones IP de forma dinámica, mediante un servidor DHCP de direcciones IPv4.

Zona 5: Se aplicó el protocolo de enrutamiento dinámico RIPv2, el cual está confeccionado con los siguientes routers: Router13, Router16, Router17 y Router21. Se crearon las siguientes VLANs denominadas: Contabilidad, Finanzas y Recursos, los equipos conectados a las VLANs obtendrán sus direcciones IP mediante un servidor DHCP de direcciones IPv4 y las VLANs están conectada al Router16.

Zona 6: Se aplicó el protocolo de enrutamiento dinámico EIGRPv2, el cual está confeccionado con los siguientes routers: Router21 y Router22. Se crearon las siguientes VLANs denominadas: Datos y VOIP. Las direcciones IP se distribuirán de forma dinámica usando DHCP en los equipos finales conectados al Router22.



Zona 7: Se aplicó un mecanismo de transición IPv4 e IPv6 conocido como Dual Stack, por lo cual toda la zona están configurado tanto en IPv4 como IPv6 y los describimos de la siguiente forma:

- En los routers Router18, Router19 y Router20 se configurara el protocolo de enrutamiento dinámico EIGRP/EIGRPv6. En el router Router19 se usa los servicios de red tales como Server HTTP, DNS y FTP para redes IPv4/IPv6.
- En los routers Router20 y Router23 se configura OSPFv2/OSPFv3, el área al que pertenecen será el Área “1”, se crearon VLANs denominadas: DatosIPv6 y DatosIPv4, las cuales estarán conectado al Router20.
- Se establece RIPv2/RIPv2 en los routers Router23, Router25, Router26 y Router27. En el Router27 existen 2 redes LAN que serán ingresadas a los equipos finales de forma manual y serán tanto IPv6 como IPv4. El Router26 ofrecerá el servicio de WIFI usando un Accesspoint inalámbrico y las direcciones IP IPv6/IPv4 serán asignadas de forma dinámica usando DHCP/DHCPv6. También algunas redes se asignaran de forma estática a los equipos finales conectados a él.

Frame Relay, NAT, NAT64, listas de acceso IPv6/IPv4 y túnel manual IPv4/IPv6.

La nube Frame Relay comunicará las redes presentes en la topología anterior. El tipo de encapsulamiento Frame Relay es cisco. Las traducciones de direcciones para IPv4 se harán de tipo NAT PAT (Port Address Translations) en el router Router13. NAT64 se implementara en el router Router2 para la traducción de direcciones IPv6 a IPv4 y viceversa. También en ese mismo router se usara en conjunto NAT64 con listas de acceso IPv6. Se establecerán 2 túneles manuales IPv4/IPv6 tal y como se muestra en la topología anterior. El túnel manual 1 estará configurado entre los routers Router2 y Router18. El túnel manual 2 por su parte se implementa en los routers Router7 y Router18.

NOTA: Recuerde que para usar NAT (PAT) se necesita hacer uso de listas de acceso IPv4. Así mismo se debe de configurar los DLCI en los routers y nube Frame Relay para que se realice la comunicación correcta en una determinada red. También se utilizará NAT64 estático y NAT-PT (IPv6) por lo cual se hace necesario el uso de listas de Acceso IPv6.

Direccionamiento IPv6 / Zona 1				
OSPFv3				
Routers	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1



Router 0	2001:6::3/64	R_Vlans	2001:5::1/64	R_Vlans
Router 4	2001:6::1/64	2001:7::1/64		
Router 5		2001:7::2/64		

Direccionamiento IPv6 / Zona 2				
EIGRPv6				
Routers	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1
Router 1	2002:1::1/64	2002:3::2/64		
Router 2	2002:5::3/64	2002:3::1/64	2002:6::1/64	2002:4::1/64
Router 3	2002:2::1/64			2002:4::2/64
Router 5	2002:5::1/64			
Router 6			2002:6::4/64	

Direccionamiento IPv6 / Zona 3				
RIPng				
Routers	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1
Router 6	2003:1::1/64			
Router 7	2003:1::2/64		2003:7::1/64	2003:2::1/64
Router 8	2003:8::1/64	R_Vlans	2003:7::4/64	2003:9::1/64
Router 9		2003:10::1/64		2003:9::2/64
Router 12		F80:DB::1/64		2003:2::2/64

Direccionamiento IPv4 / Frame Relay				
EIGRPv2				
Routers	S 0/0/0	S 0/0/1	S 0/1/1	DLCI
Router 2	192.120.70.1/28			102,103,104
Router 7		192.120.70.2/28		201,203,204
Router 13			192.120.70.3 /28	301,302,304
Router 18	192.120.70.4/28			401,402,403



Túnel					
Routers	interface Túnel	Ipv6 Address	Túnel Source	Túnel Destination	túnel Mode
Router2	Tunnel1	2010:ABCD::1/64	S 0/0/0	192.120.70.4	ipv6ip
Router18		2010:ABCD::4/64	S 0/0/0	192.120.70.1	
Router7	Tunnel2	2010:BECA::1/64	S 0/0/1	192.120.70.4	ipv6ip
Router18		2010:BECA::5/64	S 0/0/0	192.120.70.2	

Direccionamiento IPv4 / Zona 4				
OSPFv2				
Routers	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1
Router 13	172.16.40.1/30			172.16.24.1/30
Router 14	172.16.40.2/30	172.2.2.1/30		
Router 15	172.16.100.1/30	172.2.2.2/30	10.40.1.1/24	10.1.1.1/24

Direccionamiento IPv4 / Zona 5				
RIPv2				
Routers	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1
Router 13				172.16.24.1/30
Router 16	R_Vlans	172.16.8.1/30	172.16.1.1/30	
Router 17		172.16.60.1/30	172.16.1.2/30	172.16.24.2/30

Direccionamiento IPv4 / Zona 6				
EIGRP				
Routers	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1
Router 21	172.16.50.1/30	172.16.60.2/30		
Router 22	172.16.50.2/30	R_Vlans		172.18.100.1/30

Direccionamiento IPv6/IPv4 / Zona 7				
DUAL STACK				
Routers	Fa 0/0	Fa 0/1	Fa 1/0	Fa 1/1
Router 18	192.168.70.1/30			192.168.90.2/30



	2007:1::1/64			2007:3::2/64
Router 19	192.168.70.2/30 2007:1::2/64	192.168.80.1/24 2007:2::1/64		
Router 20	192.168.90.5/24 2007:5::2/64	R_Vlans		192.168.90.1/30 2007:3::1/64
Router 23	192.168.90.6/30 2007:5::1/64			192.168.90.9/30 2007:6::2/64
Router 25	192.168.100.1/30 2007:7::1/64	192.168.100.13/30 2007:9::1/64		192.168.90.10/30 2007:6::1/64
Router 26	177.16.1.1/24	192.168.100.14/30 2007:9::2/64	2007:12::1/64	177.16.2.1/24 2007:13::1/64
Router 27	192.168.100.2/30 2007:7::3/64	192.170.20.1/24	2007:8::1/64	

VLANs_Estática			
Switchs			
Nombre Switchs	Número VLANs	Nombres	Puerto Access
S0	2	Cliente	Fa 0/11-24
	3	Capacitación	Fa 0/2-10
S1	4	Audiovisual	Fa 0/11-24
	5	Conferencia	Fa 0/2-10
Switch20	100	DatosIPv4	Fa 0/12-24
	101	DatosIPv6	Fa 0/2-11

VLANs_VTP / creación			
Switch8 Switch14 Switch17			
Nombres VLANs	Número VLANs	Interfaz	Dirección Red
Oficina 1	10	Vlan 10	2003:3::/64
Oficina 2	11	Vlan 11	2003:4::/64
Oficina 3	12	Vlan 12	2003:5::/64
Oficina 4	13	Vlan 13	2003:6::/64
Contabilidad	20	Vlan 20	192.168.1.0/24



Finanza	21	Vlan 21	192.168.8.0/24
Recurso	22	Vlan 22	192.168.24.0/24
VoIP	50	Vlan 50	192.168.50.0/24
Datos	60	Vlan 60	192.168.60.0/24

VLANs_VTP / distribución		
Switches		
Nombres	Número VLANs	Puertos Access
Switch 8	10	Fa 0/20
Switch 8	11	Fa 0/10
Switch 10	12	Fa 0/5
Switch 10	13	Fa 0/11
Switch 13	20	Fa 0/2-9
Switch 13	21	Fa 0/10-18
Switch 13	22	Fa 0/19-24
Switch 14	20	Fa 0/3-10
Switch 14	21	Fa 0/11-18
Switch 17	50	Fa 0/5-15
Switch 17	60	Fa0/8, Fa 0/16-24
Switch 18	50	Fa 0/5-15
Switch 18	60	Fa0/8, Fa 0/16-24

R_VLANs				
Equipo	Interfaz	Sub-Interfaz	Dirección IP	# VLANs
R0	F 0/1	F 0/1.2	2001:1::1/64	2
		F 0/1.3	2001:2::1/64	3
	F 1/1	F 1/1.4	2001:3::1/64	4
		F 1/1.5	2001:4::1/64	5



Router8	F 0/1	F 0/1.10	2003:3::1/64	10
		F 0/1.11	2003:4::1/64	11
		F 0/1.12	2003:5::1/64	12
		F 0/1.13	2003:6::1/64	13
Router16	F 0/0	F 0/0.20	192.168.1.1/24	20
		F 0/0.21	192.168.8.1/24	21
		F 0/0.22	192.168.24.1/24	22
Router20	F 0/1	F 0/1.100	192.170.10.1/24	100
		F 0/1.101	2007:4::1/64	101
Router22	F 0/1	F 0/1.50	192.168.50.1/24	50
		F 0/1.60	192.168.60.1/24	60

Pools				
IPv6/IPv4				
Pool Name	Gateway	Dirección IP	Subredes	Inicio Dirección IP
Cliente	2001:1::1	2001:1::/48	16	2001:1::/64
Capacitación	2001:2::1	2001:2::/48	16	2001:2::/64
Audiovisual	2001:3::1	2001:3::/48	16	2001:3::/64
Conferencia	2001:4::1	2001:4::/48	16	2001:4::/64
DHCP6	2002:1::1	2002:1::/48	16	2002:1::/64
DHCP6.1	2002:2::1	2002:2::/48	16	2002:2::/64
conexipDHCPv6	2003:8::1	2003:8::/40	24	2003:8::/64
Oficina 1	2003:3::1	2003:3::/48	16	2003:3::/64
Oficina 2	2003:4::1	2003:4::/48	16	2003:4::/64
Oficina 3	2003:5::1	2003:5::/48	16	2003:5::/64
Oficina 4	2003:6::1	2003:6::/48	16	2003:6::/64



Redipv4	10.1.1.1	10.1.1.0/24		10.1.1.2/24
Redipv4.1	10.40.1.1	10.40.1.0/24		10.40.1.2/24
Contabilidad	192.168.1.1	192.168.1.0/24		192.168.1.2/24
Recurso	192.168.24.1	192.168.24.0/24		192.168.24.2/24
Finanza	192.168.8.1	192.168.8.0/24		192.168.8.2/24
Datos	192.168.60.1	192.168.60.0/24		192.168.60.2/24
Voip	192.168.50.1	192.168.50.0/24		192.168.50.2/24
Datosipv4	192.170.10.1	192.170.10.0/24		192.170.10.2/24
Datosipv6	2007:4::1	2007:4::/48	16	2007:4::/64
Wifi	177.16.2.1	177.16.2.0/24		177.16.2.2/24
	2007:13::1	2007:13::/48	16	2007:13::/64

Servidores DHCP				
Equipo	Interfaz	Subinterfaces	Dirección IP	Gateway
R0	F 0/1	F 0/1.2	2001:1::1/64	2001:1::1/64
		F 0/1.3	2001:2::1/64	2001:2::1/64
R0	F 1/1	F 1/1.4	2001:3::1/64	2001:3::1/64
		F 1/1.5	2001:4::1/64	2001:4::1/64
Router1	F 0/0		2002:1::1/64	2002:1::1/64
Router3	F 0/0		2002:2::1/64	2002:2::1/64
Router8	F 0/1	F 0/1.10	2003:3::1/64	2003:3::1/64
		F 0/1.11	2003:4::1/64	2003:4::1/64
Router8	F 0/1	F 0/1.12	2003:5::1/64	2003:5::1/64
		F0/1.13	2003:6::1/64	2003:6::1/64
	F 0/0		2003:8::1/64	2003:8::1/64
DHCP_1IPv4	F 0/0		172.16.100.2/30	172.16.100.1/30
DHCP VLAN	F 0/1		172.16.8.2/30	172.16.8.1/30
Router22	F 0/1	F 0/1.50	192.168.50.1/24	192.168.50.1/24
Server2	F 1/1		172.18.100.2/30	172.18.100.1/30
Router20	F 0/1	F 0/1.100	192.170.10.1/24	192.170.10.1/24
		F 0/1.101	2007:4::1/64	2007:4::1/64
			177.16.2.1/24	177.16.2.1/24



Router26	F 1/1		2007::13::1/64	2007::13::1/64
----------	-------	--	----------------	----------------

Dispositivo Estático				
Equipo	Zona	Dirección Red	Dirección IP	Gateway
Laptop8	1	2001:5::/64	2001:5::10	2001:5::1
Laptop16	3	2003:10::/64	2003:10::24	2003:10::1
DHCP_1IPv4	4	172.16.100.0/30	172.16.100.2	172.16.100.1
ServerHTTP	7	192.168.80.0/24	192.168.80.10	192.168.80.1
		2007:2::/64	2007:2::10	2007:2::1
ServerFTP	7	192.168.80.0/24	192.168.80.5	192.168.80.1
		2007:2::/64	2007:2::5	2007:2::1
ServerDNS	7	192.168.80.0/64	192.168.80.2	192.168.80.1
		2007:2::/64	2007:2::2	2007:2::1
DHCP VLAN	4	172.16.8.0/30	172.16.8.2	172.16.8.1
Server2	4	172.18.100.0/30	172.18.100.2	172.18.100.1

ENUNCIADO

Asignación de direcciones IPv4 e IPv6.

Se deberá asignar las direcciones IPv4/IPv6 a las interfaces FastEthernet de las PC, routers, servidores, tal y como aparecen en el cuadro llamado (**Direccionamiento IPv6/IPv4 / Zona 01-07**) y (**Dispositivo Estático**)

Creación y distribución de VLANs.

Se deberán crear y propagar las VLANs con direcciones de red IPv4/IPv6 como se muestra en las tablas llamadas: **VLANs_Estática**, **VLANs_VTP / creación** y **VLANs_VTP / distribución**.

Creación de servidores DHCPv4 y DHCPv6.

Se deberán configurar servidores DHCPv6/DHCPv4 con los pools de direcciones IP tal y como aparece en la tabla llamada **Pools**.

**Implementación de protocolos de Enrutamiento Internos en IPv4 e IPv6.**

Se deberá asignar los siguientes protocolos Internos dinámicos en IPv6: RIPng, OSPFv3 y EIGRPv6. Así mismo en IPv4 se aplicara los siguientes: RIPv2, OSPFv2 y EIGRPv2, tal y como se visualiza en la tabla denominada: **Direccionamiento IPv6/IPv4 / Zona 01-07**

Redistribución de protocolos Internos en IPv6 e IPv4.

La redistribución de rutas se configurará en el router Router5 donde se deben de anunciar el protocolo dinámico interno OSPFv3 en EIGRPv6 y viceversa que están presente en el mismo router, logrando así la comunicación entre las distintas redes que están en las zonas 1 y 2 respectivamente.

La redistribución de rutas se implementará en el router Router6 donde se anunciará el protocolo EIGRPv6 en RIPng y viceversa, proporcionando una comunicación exitosa entre las redes que están en las zonas 2 y 3.

En el router Router13 se debe de anunciar el protocolo OSPFv2 en RIPv2, EIGRPv2 en OSPFv2 y viceversa para que se lleve con éxito la comunicación entre cada una de las redes que están en las zonas 4,5 y 7 correspondientemente.

En el router Router21 se debe de anunciar el protocolo de encaminamiento EIGRPv2 en RIPv2 y viceversa para que se logre llevar a cabo la comunicación entre las redes que están en las zonas 5 y 6 de manera satisfactoria.

La redistribución de rutas se debe de establecer en el router Router18 con el protocolo de encaminamiento EIGRPv6 en RIPng y viceversa para una satisfactoria comunicación en las redes que están las zonas 3 y 7.

Se implementara la redistribución de rutas en el router Router20 donde los protocolos dinámicos EIGRP y OSPF tanto para IPv6 como IPv4 intercambiaran la información de rutas de redes que tienen en su alcance logrando que se cree la comunicación adecuada entre cada una de las diferentes redes.

Se configurará la redistribución de rutas en el router Router23 donde los protocolos dinámicos RIP y OSPF con el direccionamiento IPv6/IPv4 intercambiaran la información de rutas de redes que tienen en su alcance logrando que se cree la comunicación adecuada entre cada una de las diferentes redes existentes.

NOTA: Recuerde que algunas veces es necesario la redistribución de rutas en el mismo protocolo tal como es el caso de EIGRP en el direccionamiento IPv6.



Creación de servidores HTTP, FTP y DNS.

Los dispositivos que ofrecen estos servicios estarán en la zona 7 y su direccionamiento es IPv6/IPv4 (DUAL STACK) ya que a estos accederán muchos equipos que funcionan bajo el direccionamiento IPv6 e IPv4. Para la correcta configuración ver tabla **Dispositivo Estático**

Nota: El dominio HTTP, FTP que se registrara en el servidor DNS corresponderán con el nombre de www.tesisfinal.com y ftp.tesisfinal.com tanto para IPv6 como IPv4 respectivamente.

Configuración de Wireless con IPv6 e IPv4.

El nombre del Accesspoint es Wifi, el tipo de protección de cifrado será WPA2-PSK y su contraseña cisco123@. Las direcciones IP que ofrece a los equipos conectados a él serán IPv6/IPv4 (ambas) con el protocolo de asignación de direcciones dinámicas DHCP/DHCPv6.

Aplicación de Frame Relay, NAT64, PAT y Lista de acceso IPv4 e IPv6.

Para la configuración correcta de todos los dispositivos (routers) conectados al Frame Relay e implementación de la nube Frame Relay vea la tabla **Direccionamiento IPv4 / Frame Relay**.

El mecanismo de transición **NAT64** se implementara en el router **Router2**, en el cual también existirán listas de acceso IPv6 que se usa en combinación con el mecanismo para agrupar un conjunto de redes IPv6 en una sola dirección de red IPv4. Recuerde que en el caso de redes IPv4 se tiene que pasar una dirección exacta de la red IPv4 a IPv6 que se va a usar **NAT64** para poder comunicar estas redes entre sí.

Nota: Se crearon interfaces **Loopback** en los routers que pertenecen a las zonas 1,2 y 3 para asignar redes IPv6 a IPv4 (NAT64) para probar que se llega al menos al router que tiene la red que deseamos llegar y así tratar de resolver problemas si se presenta en las Asignación del mecanismo de transición NAT64.

PAT se implementó en el router **Router13** para brindar salida a direcciones de red privadas IPv4 que se encuentran en las zonas 4,5 y 6. También cabe destacar que pasar usar PAT (Port Address Translations) es necesario hacer uso de listas de acceso IPv4.

Implementación del mecanismo **DUAL STACK**

Se configurara el mecanismo DUAL STACK tal y como se muestra en la figura de la topología anterior. Algunos Datos estarán presente en el cuadro denominado (**Direccionamiento IPv6/IPv4 / Zona 7 /DUAL STACK**)



NOTA: Cabe mencionar que el mecanismo DUAL STACK es la implementación del direccionamiento tanto para IPv6 como IPv4 en el mismo dispositivo. Por ejemplo que el Router1 tenga configuradas 2 redes que sean IPv6/IPv4.

Implementación de Túneles Manuales

Para la correcta configuración de los túneles manuales se deben de realizar tal y como se visualiza en la tabla: **Túnel**.

NOTA: Recuerde que Aquí también se realizará distribución de rutas entre los protocolos dinámicos internos EIGRPv6, RIPng entre los routers Router2, Router7 y Router18 que usan los túneles manuales para comunicar todas las redes presentes en él.

Tiempo estimado de Solución:

- 14 horas.

Preguntas de análisis

En los siguientes apartados se pretende que los estudiantes sean capaces de analizar y responder las siguientes cuestiones en base al tema de direccionamiento IPv6/IPv4, Mecanismos de transición IPv4 e IPv6 así como el uso de protocolos de enrutamientos dinámicos internos, otros protocolos y tecnologías presente en este documento, con el fin de poner en práctica los conocimientos adquiridos tanto en la práctica como en la parte teórica.

1. En el mecanismo de transición NAT64 ¿Es necesario el uso de rutas estáticas por defecto en los routers de bordes? Explique
2. Después de haber realizado la práctica ¿En qué consiste el mecanismo de DUAL STACK?
3. ¿Explique en síntesis cual es la diferencia entre los protocolos de enrutamiento dinámico tanto en IPv6 como IPv4?
4. De Acuerdo a lo aprendido en este tema ¿Es posible la redistribución de rutas usando en el mismo protocolo? Ejemplo: EIGRPv6 en EIGRPv6 y viceversa. Justifique su respuesta.
5. En un mismo router ¿Se pueden activar listas de acceso IPv6 e IPv4?
6. ¿Cuál es la diferencia principal entre DHCPv6 y Autoconfiguración? Explique



7. Indique la afirmación que no es correcta en los siguientes enunciados:
- a) Se usa el comando IPv6-unicast para usar los protocolos dinámicos internos y externos que están desactivados por defecto.
 - b) El concepto y configuración de VLANs para IPv6 como para IPv4 es la misma.
 - c) Las redes IPv6 se comunican directamente con IPv4 y necesitan un mecanismo de transición para ello.
 - d) ninguna de las anteriores.



CAPITULO IV: ASPECTOS FINALES



1. CONCLUSIONES

Con la finalización de este trabajo monográfico, consideramos que hemos logrado cumplir con los objetivos propuestos, llegando a las siguientes conclusiones:

1. La facilitación de información teórica con ejemplos prácticos, les proporciona a los estudiantes un gran apoyo para solucionar las prácticas propuestas.
2. El diseño de un adecuado formato de prácticas facilitará a los estudiantes una correcta comprensión de la práctica a realizar.
3. La secuencia en que se han organizado las prácticas propuestas, permitirá a los estudiantes más facilidad para la solución de estas.
4. Se han abordado temas nuevos relacionados con IPv6 para el aprendizaje de los estudiantes, que son muy necesarios y han sido puesto en marcha en el ámbito laboral en el área de redes.

Con las conclusiones mencionadas, podemos afirmar que hemos logrado desarrollar un documento sencillo y práctico que los estudiantes podrán usarlo para aprender los conceptos teóricos prácticos de IPv6 y enfrentar desafíos laborales, y resolvemos las necesidades planteadas anteriormente en este documento.



2. RECOMENDACIONES

Las recomendaciones que se describen a continuación son a base de una futura actualización del documento.

1. Se ha intentado que los temas presentados aquí, estén lo más actualizados hasta la fecha de presentación de este trabajo. Sin embargo, dado que el protocolo IPv6 está siendo mejorado continuamente, es necesario que en el futuro se realice una actualización de este trabajo con las actualizaciones actuales del protocolo.
2. Es importante que, en base a este documento, se desarrolle otro trabajo con temas específicamente acerca de la Seguridad en IPv6, el cual es un tema bastante extenso y complejo, y que debido a estas características, no fue abordado en este documento.



3. BIBLIOGRAFÍA

➤ Libros digitales consultados.

- Barcha, N., Fernández, C., Frutos, S., López, G., Mengual, L., Soriano, F. J., & Yáguez, F. J. (2005). Redes de Computadores y arquitecturas de comunicaciones, Supuestos Prácticos. Madrid.
- Hogg, S., & Vyncke, E. (2009). IPv6 Security. Indianapolis: Cisco Press.
- McFarland, S., Sambi, M., Sharma, N., & Hooda, S. (2011). IPv6 for Enterprise Networks. Indianapolis: Cisco Press.
- Popoviciu, C., Levy-Abegnoli, E., & Grossetete, P. (2006). Deploying IPv6 Networks.

➤ Tesis consultadas.

- Quiroz Vázquez, R. O., Ramirez Medina, F. E., & Rivera González, Y. F. (Septiembre, 2013). Propuesta de practicas de laboratorios de switching y routing para la carrera de Ingenieria en Telemática de la UNAN-LEON. León, Nicaragua.

➤ Documentos de las web consultadas.

- Duque, S., & Vallejo, D. (s.f.). Análisis del protocolo IPv6 su Evolución y Aplicabilidad. Obtenido de <http://repositorio.utn.edu.ec/handle/123456789/1109>
- Logacho, J. E. (2014, Enero). Servicios IPv6. Obtenido de <http://dspace.ups.edu.ec/bitstream/123456789/6353/1/UPS-ST001088.pdf>
- García, J. Y. (s.f.). Movilidad en IPv6. Obtenido de http://pegaso.ls.fi.upm.es/arquitectura_redes/clase2-CUARTOIPv6MOVIL-25octubre2011.pdf
- Protocolo de Configuración Dinámica de Hosts para IPv6. (2015, Marzo 24). Obtenido de <http://es.wikipedia.org/wiki/DHCPv6>
- Narten, T., Nordmark, E., & Simpson, W. (1998, Diciembre). Neighbor Discovery for IP Version 6 (IPv6). Obtenido de [http://www.ietf.org/rfc/rfc2461.txt\\$number=2461](http://www.ietf.org/rfc/rfc2461.txt$number=2461)
- Hinden, R., & Deering, S. (2003, Abril). Internet Protocol version 6 (IPv6) Addressing Architecture. Obtenido de <http://www.ietf.org/rfc/rfc3513.txt?number=3513>



- Implementing EIGRP for IPv6. (2010, Noviembre 24). Obtenido de http://www.cisco.com/c/en/us/td/docs/ios/ipv6/configuration/guide/12_4t/ipv6_12_4t_book/ip6-eigrp.html
- Implementing IPsec in IPv6 Security. (2011, marzo 25). Obtenido de http://www.cisco.com/c/en/us/td/docs/ios/ipv6/configuration/guide/12_4t/ipv6_12_4t_book/ip6-ipsec.html
- Implementing Multiprotocol BGP for IPv6. (2011, septiembre 26). Obtenido de http://www.cisco.com/c/en/us/td/docs/ios/ipv6/configuration/guide/12_4t/ipv6_12_4t_book/ip6-mptcl_bgp.html
- Implementing OSPF for IPv6. (2011, Septiembre 26). Obtenido de http://www.cisco.com/c/en/us/td/docs/ios/ipv6/configuration/guide/12_4t/ipv6_12_4t_book/ip6-ospf.html
- Implementing RIP for IPv6. (2011, Septiembre 26). Obtenido de http://www.cisco.com/c/en/us/td/docs/ios/ipv6/configuration/guide/12_4t/ipv6_12_4t_book/ip6-rip.html
- Implementing Static Routes for IPv6. (2011, Julio 28). Obtenido de http://www.cisco.com/c/en/us/td/docs/ios/ipv6/configuration/guide/12_4t/ipv6_12_4t_book/ip6-stat_routes.html
- Implementing Tunneling for IPv6. (2011, Septiembre 26). Obtenido de http://www.cisco.com/c/en/us/td/docs/ios/ipv6/configuration/guide/12_4t/ipv6_12_4t_book/ip6-tunnel.html
- Implementing VoIP for IPv6. (2011, Julio 22). Obtenido de http://www.cisco.com/c/en/us/td/docs/ios/ipv6/configuration/guide/12_4t/ipv6_12_4t_book/ip6_voip.html
- Davies, J. (2007, Julio). Tráfico IPv6 sobre conexiones VPN. Obtenido de <https://technet.microsoft.com/es-es/magazine/2007.07.cableguy.aspx?pr=PuzzleAnswer>